

INTEGRATING SUPERCONDUCTING QUBITS WITH QUANTUM FLUIDS AND
SURFACE ACOUSTIC WAVE DEVICES

By

Justin R. Lane

A DISSERTATION

Submitted to
Michigan State University
in partial fulfillment of the requirements
for the degree of

Physics - Doctor of Philosophy

2021

ABSTRACT

INTEGRATING SUPERCONDUCTING QUBITS WITH QUANTUM FLUIDS AND SURFACE ACOUSTIC WAVE DEVICES

By

Justin R. Lane

Superconducting qubits, mesoscopic superconducting circuits with a single quantum coherent degree of freedom, have emerged as both a promising platform for quantum computation and a versatile tool for creating hybrid systems with quantum mechanical degrees of freedom. In this dissertation, we report on several experiments investigating the coupling of superconducting 3D transmon qubits to two different systems: superfluid helium and piezoelectrically actuated surface acoustic waves (SAWs). We report the first measurements of superconducting qubits in the presence of superfluid helium, studying the spectroscopic and decoherence properties of this combined system. We analyze the spectroscopic properties of this composite system using the framework of circuit quantum electrodynamics, and in the presence of superfluid helium we observe modest increases in the pure dephasing time. We attribute this to improved thermalization of the microwave environment via the superfluid, raising hopes that thermalization mediated by superfluid helium may be a resource for experiments employing superconducting circuits. We also present ongoing work developing a new capacitive coupling scheme for creating hybrid superconducting qubit-SAW resonators. The tools developed to model and implement this experiment lay the groundwork for future experiments to achieve robust coupling between 3D transmon qubits and surface acoustic wave devices. Finally, we present results describing the first measurements of SAW induced transport in exfoliated graphene devices.

ACKNOWLEDGMENTS

I am of the conviction that it takes a village to raise a scientist. Because of that, I cannot possibly thank everyone who contributed to my Ph.D. in one way or another. However try my best to acknowledge everyone who was pivotal to the process.

First off, I'd like to thank the person who was most important to this process: my advisor, Johannes Pollanen. Johannes is an amazing scientist and role model, and I have been incredibly lucky to have had his guidance for the past 6 years. Johannes has given me more trust and academic freedom than a grad student could reasonably ask for, and has always been available and willing to lend help when I need it (which is a lot.) He has been an excellent teacher of both science and the process of doing science, and yet somehow has always made me feel like a colleague rather than a student. I would be nowhere near the scientist I am today without his guidance, patience, and friendship, and can only hope to one day be as good a mentor to someone else.

In addition to Johannes, would also like to thank the members of my guidance committee: Prof. Dean Lee, Prof. Chris Wrede, Prof. Norman Birge and Prof. Mark Dykman. I am particularly in debt to Mark and Norman, who have been both willing to share a great deal of knowledge and expertise with me and excellent at doing so.

On top of the senior scientists at MSU who have guided me for the past several years, I am lucky to have had the opportunity to collaborate with a group of excellent physicists outside of MSU. Chief among those I have to thank is Prof. Kater Murch at Washington University in St. Louis, without whom this dissertation would not have been possible. Kater's expertise in superconducting qubits, and his willingness to share that expertise, has been absolutely foundational to the work presented here. I also must thank Prof. Erik Henriksen, also at

Wash. U., who played a pivotal role early in my Ph.D. in guiding the graphene experiments that appear in this thesis. Thank you also to the grad students at Wash. U. who I have had the pleasure to collaborate with over the years: Dr. Dian Tan, Daria Kowsari, and Dr. B. “Nero” Zhou.

I have been fortunate to have the opportunity to work with many amazing scientists in LHQS over the past 6 years. I am indebted to all of the postdocs who I have had the opportunity to work with, and who have lent me much of the guidance that has made this dissertation possible: Dr. Mazin Khasawneh, Dr. Kostya Nasyedkin, and especially Dr. Niyaz Beysengulov. Special thanks must also be given to the original LHQS crew, Liangji Zhang and Dr. Heejun Byeon, who I’ve been lucky to call colleagues and friends ever since the lab was an empty room.

Some of the best advice I have ever been given is that folks both younger and smarter than you *will* appear in your life, and rather than being bitter about it you should take every opportunity to learn as much as possible from them. With that in mind, I am indebted to Joe Kitzman, Camille Mikolas, and Steve Hemmerle, who have made the last several years both an excellent learning experience and a lot of fun. I would also like to thank the small army of undergraduate students who have passed through LHQS in my time: Anna Turnbull, Josh Milem, Taryn Stefanski, Evan Brook, Gabe Moreau, and Brennan Arnold.

There are many other people in the PA department for whom I am better off having known. I’d like to thank co-organizers of the QuIC seminar, who out of shared interest and determination made the seminar a reality: Ryan LaRose, Ben Hall, Niyaz Beysengulov (again) and Jacob Watkins. I am also indebted to many members of the department support staff, who have been an invaluable source of help during my time at MSU. Special thanks to Dr. Reza Loloei, Dr. Baokang Bi, Kim Crosslan, Cathy and Jessica Cords, and to the shop

guys: Tom Palazzolo, Jim Muns, Tom Hudson and Rob Bennett.

I, like roughly everyone else on Earth, had my life turned upside-down by the pandemic, and have been lucky enough to have a great group of friends who have been a source of comfort from afar during this past year: thank you to Eli Clausen, Travis Schess, Gabe Wascher, Drew Jordan, Zach Campbell, and Nick and Kelly Jaroz for your support and friendship from a distance during this past, fairly arduous year. Of all of these folks, I am particularly indebted to Nick, who 9 years ago was the first person to convince me that a normal person such as myself could build a career out of intellectual curiosity. I would also like to thank Rob Elder and Kolten LaBarre, who in addition to being great housemates made the first and most uncertain months of the pandemic much more enjoyable.

I am deeply indebted to my family: my parents Monica and Rich Lane, my sister Meghan and her husband Joe Carlson. Without your support for the past 6 (and, frankly, 28) years, none of this would have been possible. Finally, I would like thank Kailey Draves, who has put up with me more than any other person during this process. You're ok kid.

The projects described in this dissertation were partially supported by the National Science Foundation under grant numbers DMR-1708331, DMR-2003815 and DMR-1810305, and by the generous support of the Jerry Cowen endowment.

TABLE OF CONTENTS

LIST OF TABLES	ix
LIST OF FIGURES	x
Chapter 1 Introduction	1
1.1 More is different (and surprisingly still quantum)	2
1.2 Quantum acoustics	3
1.3 A short outline of this dissertation	5
Chapter 2 Superconducting circuits and circuit QED	6
2.1 A brief history of superconducting quantum circuits	7
2.2 The building blocks of quantum circuits	9
2.2.1 Quantizing the LC oscillator	10
2.2.2 Josephson junctions	14
2.2.3 The Cooper pair box	17
2.2.4 Energy scales	21
2.2.5 The transmon	22
2.2.6 Where did the charge noise go?	27
2.3 Circuit quantum electrodynamics	29
2.3.1 Putting together the building blocks of cQED	30
2.3.1.1 Driving the system	30
2.3.1.2 Coupling a qubit and a harmonic oscillator	32
2.3.2 The Jaynes-Cummings model	34
2.3.3 The strong dispersive limit and readout in cQED	37
Chapter 3 Superconducting qubits in the lab	40
3.1 The 3D transmon geometry	40
3.1.1 Electromagnetic modes in 3D cavities	41
3.1.2 3D transmon qubits	44
3.2 Filtering and shielding	49
3.2.1 Cryogenic shielding	50
3.2.2 Input/output line filtering	51
3.2.3 Amplification	54
3.2.4 Possible improvements	56
3.3 Spectroscopic measurements	59
3.3.1 “Punchout”	59
3.3.2 Two-tone spectroscopy	61
3.4 Time-resolved measurements	63
3.4.1 General time-resolved setup	63
3.4.2 Coherent control	64
3.4.3 Punchout-based readout	69

3.4.4	T_1, T_2 and T_2^e	72
Chapter 4	Integrating superfluids with superconducting qubit systems	77
4.1	Superfluid helium	79
4.2	Decoherence and noise	81
4.2.1	Common noise sources in superconducting qubits	85
4.3	Experimental setup	88
4.3.1	Gas handling system and fill line	88
4.3.2	Superfluid helium leak tight sample cell	92
4.4	Experimental results	94
4.4.1	Cavity and qubit spectroscopy	94
4.4.2	Qubit Relaxation and Decoherence	100
4.4.2.1	Qubit energy relaxation	102
4.4.2.2	Qubit dephasing	105
4.4.2.3	Long timescale fluctuations in qubit coherence properties	107
4.4.3	Residual excited state population	109
4.5	Discussion and future work	110
4.5.1	Coupling to superfluid acoustic modes?	114
Chapter 5	Surface acoustic waves and surface acoustic wave devices	117
5.1	Elastic waves in solids	118
5.1.1	Rayleigh Waves	123
5.2	Piezoelectricity and piezoelectric Rayleigh waves	129
5.2.1	Equations of motion and Boundary conditions	130
5.2.2	Sketch of Piezoelectric Rayleigh waves	132
5.2.3	Relevant parameters	134
5.3	SAW devices	136
5.3.1	Interdigitated transducers	136
5.3.2	Coupling of Modes and the P-matrix	140
5.3.3	A second look at transducers	143
5.3.4	Acoustic Bragg mirrors	144
5.3.5	SAW resonators	145
Chapter 6	Quantum acoustics using surface acoustic waves	150
6.1	Mediating the coupling between SAWs and qubits	152
6.2	Device design	154
6.2.1	Modeling the experiment	156
6.2.1.1	Finite element modeling	157
6.2.1.2	Extracting the Butterworth-van Dyke equivalent circuit	160
6.2.1.3	Classical circuit simulations	162
6.2.2	Optimizing the SAW device parameters	165
6.3	Device fabrication and assembly	166
6.4	Experimental results	169
6.4.1	Devices on ST-X quartz	169
6.4.1.1	Uncontrolled substrate spacing	173

6.4.1.2	Mirror reflectivity	175
6.4.1.3	IDT position between the mirrors	176
6.4.1.4	Coherence times	178
6.4.2	Devices on LiNbO ₃	179
6.5	Future work	181
Chapter 7	Acoustoelectric transport in graphene	183
7.1	Graphene: a brief introduction	184
7.2	Surface acoustic waves and two-dimensional electron systems	186
7.2.1	SAW attenuation and velocity shift	187
7.2.2	Acoustoelectric effect	190
7.3	Acoustoelectric effect in exfoliated graphene	193
7.3.1	Cold-finger and sample stage	194
7.3.2	Experimental setup	198
7.3.3	SAW delay line characterization	201
7.3.4	Graphene acoustoelectrics	202
7.4	Future work	207
7.4.1	Graphene acoustoelectrics in a magnetic field	207
7.4.2	$\Delta v/v$ measurements	208
7.4.3	Measurements of 2DEs using SAW resonators	209
APPENDICES		211
APPENDIX A	A note on the rotating wave approximation	212
APPENDIX B	Other cQED experimental techniques	215
APPENDIX C	Fabrication Recipes	222
BIBLIOGRAPHY		236

LIST OF TABLES

Table 3.1: RF circuitry components for the circuit quantum electrodynamics experimental setup, along with part numbers and suppliers. For some components, we use multiple variants of the same component. For these cases, the variable parameter is marked in bold , and the corresponding change in the part number is matched to the variable in the component description. Some of the specific products here are no longer manufactured, however as of writing this, every manufacturer listed either still makes the part listed or makes an updated version of it.	57
Table 4.1: Spectroscopic parameters of the cavity/qubit system both in the presence and absence of superfluid helium. $\omega_c, \delta\omega, \omega_{01}$, and ω_{12} are measured values, while E_J, E_C and g_{01} are extracted by solving the generalized Jaynes-Cummings Hamiltonian constrained by measured spectroscopic parameters.	96
Table 5.1: Relevant material parameters for common SAW substrates, reproduced from Chapter 4 of Ref. [160] and Chapter 9 of Ref. [167], both of which contain more complete tables of materials/parameters.	135
Table 6.1: Summary of the design parameters, phenomenological parameters, and derived parameters used in the model of the example SAW resonator considered in § 6.2. All design parameters are defined in fabrication.	164
Table 6.2: Summary of the target design parameters, phenomenological parameters, and derived parameters for the device on ST-X quartz experimentally tested and described in § 6.4.1.	170
Table 6.3: Transmission characteristics of a 3D cavity with/without LiNbO ₃ , extracted from the data plotted in Fig. 6.13	180

LIST OF FIGURES

Figure 2.1: A simple LC circuit	10
Figure 2.2: The Josephson junction and the Cooper pair box	15
Figure 2.3: Eigenvalues of the Cooper pair box Hamiltonian	20
Figure 2.4: A mechanical analogue to the Cooper pair box Hamiltonian	24
Figure 2.5: Schematic of circuit QED	31
Figure 2.6: Cavity-qubit avoided crossing	37
Figure 2.7: Cavity dispersive shift	38
Figure 3.1: 3D electromagnetic cavities	43
Figure 3.2: Cavity decay rate κ vs. pin depth	45
Figure 3.3: The 3D transmon qubit geometry	46
Figure 3.4: Cryogenic shielding for superconducting qubit experiments	50
Figure 3.5: Filtering and amplification of the input-output lines	53
Figure 3.6: Qubit “punchout”	60
Figure 3.7: Two-tone spectroscopy	62
Figure 3.8: General time resolve measurement setup	64
Figure 3.9: Rabi oscillations	68
Figure 3.10: IQ mixer modulation and demodulation	71
Figure 3.11: IQ plane measurement and binning	72
Figure 3.12: Measuring T_1 , T_2 , and T_2^e	74
Figure 4.1: Noise and decoherence of a qubit	82
Figure 4.2: Gas handling system and cold trap	89

Figure 4.3: Superfluid Helium + superconducting qubit setup	91
Figure 4.4: Superfluid leak-tight sample cell, iteration #1	93
Figure 4.5: Cavity and qubit spectroscopy in the presence of superfluid helium	95
Figure 4.6: Simplified circuit model for finding Δg_{01}	98
Figure 4.7: Description of time-resolved and averaged measurements	101
Figure 4.8: T_1 , T_2 and ω_{01} as a function of temperature with/without superfluid helium	102
Figure 4.9: Long timescale fluctuations of coherence parameters with/without superfluid helium	107
Figure 4.10: Excited state population of the qubit as a function of temperature with/without superfluid helium	109
Figure 4.11: Proposals for future superfluid-superconducting qubit experiments	115
Figure 5.1: Stress and strain in elastic media	119
Figure 5.2: SAW propagation in an isotropic medium	127
Figure 5.3: Interdigitated transducers for SAWs	137
Figure 5.4: The P-matrix for SAW devices	141
Figure 5.5: Effect of reflections on IDT conductance	143
Figure 5.6: Bragg Mirrors	146
Figure 5.7: COM model of a SAW resonator	147
Figure 6.1: Schematic of the proposed experiment	153
Figure 6.2: Finite element simulations of the paddle structure	158
Figure 6.3: Butterworth-van Dyke equivalent response of a SAW resonator	160
Figure 6.4: Extracting the Hamiltonian parameters from classical circuit simulations	162
Figure 6.5: Finite element simulations of the paddle structure	164
Figure 6.6: Mask for aligning the flip-chip device stack	167

Figure 6.7: Assembling the flip-chip stack	168
Figure 6.8: Verifying coupling of the transmon-like mode to the electromagnetic cavity	171
Figure 6.9: Spectroscopy of the ST-X Quartz quantum acoustics device	172
Figure 6.10: Dependence on the substrate spacing z	174
Figure 6.11: Dependence of the SAW response on the grating reflectivity	175
Figure 6.12: Dependence of g_m on the IDT position	177
Figure 6.13: 3D cavity transmission in the presence of LiNbO_3	180
Figure 7.1: DC transport measurements of graphene	186
Figure 7.2: SAW attenuation and velocity shift vs. 2DES conductivity	189
Figure 7.3: Cold-finger design	195
Figure 7.4: Acoustoelectrics sample stage	197
Figure 7.5: Graphene acoustoelectrics device schematic	199
Figure 7.6: SAW reflection and acoustoelectric voltage vs. drive frequency	202
Figure 7.7: Acoustoelectric effect: charge carrier polarity dependence	204
Figure 7.8: Gate-tunable acoustoelectric effect	205
Figure 7.9: Power dependence of the Gate-tunable acoustoelectric effect	206
Figure A.1: The rotating wave approximation	213
Figure B.1: Excited state population measurement	216
Figure B.2: Spin-locking setup and measurement	219
Figure B.3: Spin-locking spectroscopy	220
Figure C.1: Description of the Dolan bridge junction fabrication process	229
Figure C.2: SAW device fabrication	231
Figure C.3: Examples of SAW IDT dosing effects	233

Chapter 1

Introduction

Nearly a century after Schrödinger wrote down his famous equation, largely solidifying the formalism of quantum mechanics still in use today, the fact remains that we live in a *classical* world. Everyone agrees that quantum mechanics is living under the hood of every physical system, but at the scale of us humans it is completely washed out. This suppression of quantum physics is nearly complete: you don't need to understand quantum mechanics to play baseball, or to build a car, or even to put a person on the moon, even though all of these activities are centered around the motion of objects whose constituents are supposed to be quantum mechanical. A pitcher need not worry about the Heisenberg uncertainty relationship when throwing a curve ball, an engineer doesn't consider entanglement when designing a crank shaft, and Houston didn't consider the possibility of the lunar lander in a superposition of landing on the Moon and Mars.

The emergence of our classical, macroscopic world from the underlying “quantum-ness” of its constituent particles remains an important question in fundamental physics and active area of research. Conversely, one may ask how far can we push the concept of “quantum-ness”? Can we engineer a “macroscopic” system that obeys the laws of quantum mechanics? How far can we depart from the realm of single particles and still maintain an object that obeys the Schrödinger equation?

1.1 More is different (and surprisingly still quantum)

Phil Anderson, in his landmark 1972 paper [1], coined the term “more is different” to announce the idea that the study of emergent complexity, rather than simply being an application of the laws governing the system’s constituents, is itself a fundamental scientific enterprise. Anderson’s main thesis is that reductionism does not imply constructionism; that the emergent degrees of freedom of a complex system cannot, in general, be *a priori* predicted from the laws governing the underlying constituents. In the 50 years since, “more is different” has become somewhat of a mantra of condensed matter physics, which can be summarized as the study of the emergent degrees of freedom in many-body systems whose constituents are (usually) electrons and ions interacting electromagnetically.

But is there any reason to believe that these many-body degrees of freedom should behave quantum mechanically? For example, the vibrations of ions in a crystal may be quantized into phonons, which may be thought of as “particles” in the same way photons are quanta of the electromagnetic field. While phonons are invaluable in understanding many properties of condensed matter systems, is it valid to think of them as quantum mechanical particles, like photons? Indeed, our experience tells us that sufficiently large many-body systems, baseballs and crankshafts and spaceships, tend to behave classically. Can the collective motion of trillions of atoms in a crystal be in a superposition state? Can a many-body degree of freedom be entangled with another particle?

The answer to this question is, perhaps surprisingly, a resounding *yes*. In a groundbreaking experiment in 2010, O’Connell et al. [2] demonstrated coherent coupling of a dilatational mode of a microelectromechanical (MEMS) device to a superconducting quantum bit, or qubit. As we shall see in Chapters 2 and 3, superconducting qubits are mesoscopic

circuits (typically $\sim 1\ \mu\text{m} - 1\ \text{mm}$) with a quantum coherent degree of freedom defined by the rigid motion of the superconducting condensate¹. Originally designed to be the fundamental computing element of a quantum computer, these devices may be engineered into a wide variety of geometries, with fine control of the associated electric field profile. O’Connell et al. used this flexibility to create a superconducting qubit device that coupled strongly to a MEMS resonator via the piezoelectric effect. The authors were able to demonstrate coherent exchange of a single quanta between the qubit and the dilatational mode of the MEMS resonator, and even measure the phase coherence of a superposition state of the resonator, an unambiguous verification of the quantum nature of the collective mode.

1.2 Quantum acoustics

In the years since, the model of leveraging the quantum coherence of superconducting circuits to study and manipulate acoustic degrees of freedom at the single quantum level has blossomed into a research program in its own right. Superconducting qubits have been employed to control several different types of acoustic devices at the single quantum level, including MEMS structures [2, 3, 4], bulk acoustic wave resonators [5, 6], and surface acoustic wave devices [7, 8, 9, 10, 11]. Many impressive experimental results have been demonstrated, including single phonon splitting of the qubit spectrum [4, 12], Wigner function negativity of the acoustic resonator [10, 13], squeezed states of sound [14] and erasure of which-path information of a propagating phonon [15], demonstrating unambiguously the quantum nature of these collective excitations.

¹Superconducting qubits are themselves a profound answer to the question of how large can we make a system that behaves like a single quantum degree of freedom. As we’ll see, the answer turns out to be “visible to the naked eye.”

This field of research has been branded “quantum acoustics” in analogy to its electromagnetic parallel quantum optics. In quantum optics, we typically consider the interaction of a quantum system (such as an atom) with electromagnetic radiation; from this perspective, quantum *acoustics* is the interaction of a single quantum system (a superconducting qubit) with acoustic radiation. The fundamental difference between quantum acoustics and quantum optics is that the speed of sound is $10^4 - 10^5$ times slower than the speed of light (meaning that the wavelength of sound at a given frequency is $10^4 - 10^5$ times smaller than corresponding electromagnetic radiation.) This disparity in time/length opens up regimes of study that aren’t accessible in traditional quantum optics: for example, one may enter the “giant atom” regime, where the size of the single quantum object (the qubit) is much larger than the wavelength of the radiation it interacts with [16, 17].

Not too long ago, the question of engineering a many-body quantum system was purely academic. However, in parallel with the rise of experimental quantum information processing systems and nascent quantum computers, physicists have already started to ask how quantum acoustics may be leveraged in these applications. Again, the low propagation speed/small wavelength are the primary resource of quantum acoustics. $\sim$ GHz frequency acoustic resonators may be fabricated with high quality factors and small spatial extent, raising the possibility of dense bosonic quantum memories [18]. The wavelength of sound at telecom frequencies (where qubits are typically operated) may also be engineered to match common optical wavelengths, opening up an avenue to coherent microwave-to-optical conversion [19].

1.3 A short outline of this dissertation

This dissertation represents, for lack of a better term, a grab bag of experiments, motivated by the recent progress in quantum acoustics, that lays the foundation of a similar research program at MSU. We will begin with a theoretical description of superconducting qubits in Chapter 2, followed by a survey of the experimental capabilities of our lab in Chapter 3. In Chapter 4, we will report on experiments investigating the compatibility of superconducting circuits with another system known to host many interesting mechanical excitations: superfluid helium. We will then spend Chapters 5 and 6 discussing surface acoustic waves (SAWs), how we might leverage our experimental capabilities to create novel geometries for quantum acoustics with SAWs, and finally report on preliminary experimental results along this front. Chapter 7 catalogues an experiment from earlier in my Ph.D. using SAWs to study an entirely different many-body system: graphene. While this experiment doesn't directly relate to superconducting qubits and quantum acoustics, much of the experimental knowledge of SAW techniques in our lab was developed during this experiment. I also believe there is ample room for overlap between SAW experiments studying low dimensional electron systems and many other ideas outlined in this thesis, which I will attempt to motivate in the future work section of Chapter 7.

Chapter 2

Superconducting circuits and circuit

QED

The extraordinary rise of computing power over the past 50 years is a direct product of the massive scalability and miniaturizability of complementary metal oxide silicon (CMOS) integrated circuits. The top-down fabrication of metal-oxide transistors via photolithography has allowed for device complexity to scale more or less exponentially, to the point where a high-end phone today has on the order of 10 billion transistors.

Ever since the infancy of quantum computing, the search for a physical platform on which to perform quantum computations has been predicated on this requirement of massive scalability. Naturally, the success rendered by the top-down fabrication of classical integrated circuits has often percolated this discussion: is it possible to process quantum information using circuits fabricated with the same photolithographic techniques used in CMOS integrated circuits? From this line of thought, superconducting circuits have emerged as a promising candidate for quantum computing. In this chapter, we survey several fundamental elements of superconducting circuits, and discuss the circuit quantum electrodynamics (cQED) architecture for manipulating and reading out qubits based on superconducting circuits.

2.1 A brief history of superconducting quantum circuits

In 1911, H. Kamerlingh Onnes discovered that, upon cooling down to liquid helium temperatures, the electrical resistance of mercury completely disappeared [20]. Since then, many metals have been found to transition into a superconducting state below some critical temperature T_c , characterized by no dc resistance and complete expulsion of external magnetic fields [21]. While the study of superconductivity has been active for over 100 years, this brief discussion will be constrained to several theoretical and experimental breakthroughs important to the realization of superconducting quantum circuits.

In 1950, Ginzburg and Landau (GL) proposed a phenomenological theory of superconductivity, where the superconducting state is characterized by a complex order parameter $\psi(\mathbf{r})$, interpreted as a pseudo-wave function describing the superconducting state [22]. Seven years later, Bardeen, Cooper and Schrieffer (BCS) proposed a fully microscopic theory where paired electrons, called Cooper pairs, condense into a macroscopic ground state [23]. One hallmark of the BCS theory was the prediction that single particle excitations above the ground state are gapped, with an energy gap Δ_{BCS} comparable to the thermal energy at the superconducting transition temperature $k_B T_c$. This gap parameter, which in general may be a complex function of the crystal momentum $\mathbf{k}$ (and thus $\mathbf{r}$) was soon shown [24] to be directly proportional to the GL order parameter $\psi(\mathbf{r})$, justifying the phenomenological nature of the GL theory.

In 1962, Josephson observed [25] that when two superconductors are connected by a tunnel barrier, there may be a discontinuity in the phase of the complex order parameter across the barrier. This observation led to prediction of the so called Josephson effect, which

was shortly thereafter observed by Anderson and Rowell [26]. The current and voltage across these tunnel barriers, called Josephson junctions, may be written in terms of the phase difference $\varphi = \varphi_2 - \varphi_1$:

$$I = I_c \sin \varphi \tag{2.1a}$$

$$\frac{d\varphi}{dt} = \frac{2eV}{\hbar} \tag{2.1b}$$

Here, I_c is the critical current of the Josephson junction: the maximum amount of supercurrent that may pass through the junction before the onset of dissipative current. Josephson junctions essentially represent a new superconducting lumped circuit element with a nonlinear current-voltage relationship. The nonlinearity of the Josephson junction forms the backbone upon which superconducting quantum circuits are built.

While superconductivity is generally described as a macroscopic quantum phenomenon, in 1980 Leggett argued that the macroscopic effects of superconductivity, such as the Josephson effect, can be viewed as arising from the product of many single Cooper pair (two particle) wave functions rather than a coherent superposition of one many-body wave function [27]. In the same manuscript, Leggett proposed that truly macroscopic quantum effects may be observed in a superconducting ring interrupted by a Josephson junction (an RF superconducting quantum interference device, or SQUID) by leveraging the nonlinearity of the Josephson junction. It is possible to write the potential energy of a SQUID loop as a function of the phase difference across a junction, and to have a potential landscape where more than one local minimum in the phase variable may exist. It was argued that such a “phase particle”, a macroscopic state of many Cooper pairs, may be able to undergo quantum tunneling

between two potential wells, an inherently quantum phenomenon.

This prediction was followed up by a series of seminal experiments by Clarke et al. showing that the phase difference across a Josephson junction could, in fact, act as a quantum “phase particle” in a potential well, with quantized energy levels and a finite tunneling probability into an adjacent well [28]. This experiment, and many that followed, opened up the possibility of using Josephson junction-based superconducting circuits to engineer and manipulate macroscopic quantum-mechanical degrees of freedom. Coherent manipulation of a macroscopic quantum variable in a Josephson junction circuit was demonstrated by Nakamura et al. [29], ushering in the era of superconducting qubits, where the many-body quantum states of superconducting circuits could be coherently manipulated and read out.

2.2 The building blocks of quantum circuits

The analysis of superconducting quantum circuits necessitates the combination of quantum mechanics (generally written in the language of Hamiltonians and operators) and circuit analysis (generally written in terms of currents and voltages). Combining these languages is somewhat counterintuitive at first, so before we discuss Josephson junction-based qubits and superconducting circuits it will be useful to build intuition by considering a more familiar case: a lossless LC circuit. Much of the following discussion is based in piecemeal fashion off Ref. [30] and chapter 3 of Ref. [31]. A detailed discussion on circuit quantization may be found in Ref. [32].

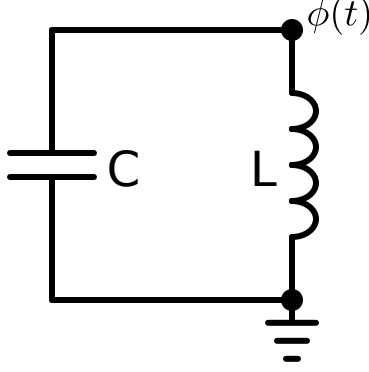


Figure 2.1: **A simple LC circuit:** A parallel LC circuit. We have defined two nodes: one at ground, and one at node flux relative to ground $\phi(t)$

2.2.1 Quantizing the LC oscillator

Consider a parallel LC circuit, like the one shown in Fig. 2.1. For reasons that will become apparent shortly, instead of describing the circuit in terms of voltages and currents, we will describe the circuit in terms of the *node flux* $\phi(t)$ of the nodes (i.e., wires connecting elements) of the circuit

$$\phi(t) = \int_{-\infty}^t V(t') dt' \quad (2.2)$$

Here, $V(t')$ is the voltage of the node relative to ground: we are explicitly constructing $\phi(t)$ such that $d\phi(t)/dt = \dot{\phi} = V(t)$. For the parallel LC circuit in Fig. 2.1, we have two nodes: one at ground potential, and one node “above” ground potential at $\phi(t)$. From elementary circuit analysis, we know that the voltage across the inductor is $V = L\dot{I} = \dot{\phi} \rightarrow LI = \phi$ for a static inductance, where L and I are the inductance and current across the inductor respectively. We also know that the energy stored in the inductor and the capacitor may be written as

$$\begin{aligned}
U_C &= \frac{1}{2}CV^2 = \frac{1}{2}C\dot{\phi}^2 \\
U_L &= \frac{1}{2}LI^2 = \frac{1}{2L}\phi^2
\end{aligned}
\tag{2.3}$$

From this point of view, the energy in the capacitor looks like the *kinetic* energy of a particle with position ϕ , and the energy in the inductor looks like the *potential* energy. Keeping this convention in mind, we write down the Lagrangian of the system and find the conjugate momentum Q to our coordinate ϕ :

$$\mathcal{L} = U_C - U_L = \frac{1}{2}C\dot{\phi}^2 - \frac{1}{2L}\phi^2 \tag{2.4}$$

$$Q = \frac{\partial \mathcal{L}}{\partial \dot{\phi}} = \dot{\phi}C = VC \tag{2.5}$$

which is, of course, the charge on the capacitor plates. Our convenient choice of ϕ as a coordinate has given us a set of physically meaningful conjugate variables! We can then Legendre transform the Lagrangian and write down the Hamiltonian for the circuit in terms of the two conjugate variable ϕ and Q :

$$H = Q\dot{\phi} - \mathcal{L} = \frac{Q^2}{2C} + \frac{\phi^2}{2L} \tag{2.6}$$

which, since there is no explicit time dependence, is simply the sum of the energies U_C and U_L . With the system Hamiltonian written in terms of a set of conjugate variables, we may now safely promote the Hamiltonian/variables to operators with a canonical commutation relationship $[\phi, Q] = i\hbar$. If we recall that the resonant frequency of a simple LC oscillator is $\omega_0 = (LC)^{-1/2}$, the fully quantum mechanical Hamiltonian may be written as

$$\mathbf{H} = \frac{1}{2C}\mathbf{Q}^2 + \frac{1}{2}C\omega_0^2\phi^2 \quad (2.7)$$

We recognize this as exactly the Hamiltonian of a 1D harmonic oscillator, with the mass replaced by the capacitance, and $\mathbf{x}$ and $\mathbf{p}$ replaced with ϕ and $\mathbf{Q}$ respectively. Following the standard treatment of a quantum harmonic oscillator, we construct a set of raising and lowering operators for our circuit:

$$\phi = \phi_{ZPF}(\mathbf{a} + \mathbf{a}^\dagger) \quad (2.8a)$$

$$\mathbf{Q} = -iQ_{ZPF}(\mathbf{a} - \mathbf{a}^\dagger) \quad (2.8b)$$

$$[\mathbf{a}^\dagger, \mathbf{a}] = 1 \quad (2.8c)$$

We can then write down the Hamiltonian in terms of these raising and lowering operators

$$\mathbf{H} = \hbar\omega_0\left(\mathbf{a}^\dagger\mathbf{a} + \frac{1}{2}\right) \quad (2.9)$$

with the familiar result for the energy $E = \hbar\omega_0(N + 1/2)$. The prefactors to the right hand sides of Eqns 2.8a and 2.8b are the magnitudes of the zero point fluctuations of ϕ and $\mathbf{Q}$ respectively, and are given by:

$$\phi_{ZPF} = \sqrt{\frac{\hbar Z_0}{2}} \quad (2.10a)$$

$$Q_{ZPF} = \sqrt{\frac{\hbar}{2Z_0}} \quad (2.10b)$$

$$Z_0 = \sqrt{\frac{L}{C}} \quad (2.10c)$$

where Z_0 is the characteristic impedance of the oscillator. Why do we call them the zero point fluctuation magnitudes? In the ground state of the harmonic oscillator it is easy to show that

$$\begin{aligned} \langle \phi \rangle &= \phi_{ZPF} \langle 0 | (\mathbf{a}^\dagger + \mathbf{a}) | 0 \rangle = 0 \\ \langle \phi^2 \rangle &= \phi_{ZPF}^2 \langle 0 | (\mathbf{a}^{\dagger 2} + \mathbf{a}^2 + \mathbf{a}\mathbf{a}^\dagger + \mathbf{a}^\dagger\mathbf{a}) | 0 \rangle = \phi_{ZPF}^2 \end{aligned} \quad (2.11)$$

We therefore see that the zero point variance of ϕ is $\Delta\phi^2 = \langle \phi^2 \rangle - \langle \phi \rangle^2 = \phi_{ZPF}^2$. The same may be shown for $\mathbf{Q}$.

Wait a minute. Is this valid? Compared to what we normally consider “quantum” scale systems, capacitors and inductors are *macroscopic* objects: even a $\sim \mu\text{m}$ scale piece of metal has order 10^{12} electrons. How can we talk about a *single* quantum degree of freedom arising from this hopelessly many-body system?

The answer arises from the physics of superconductors: in a superconductor, single particle excitations above the superconducting ground state are *gapped* by an energy Δ_{BCS} . The existence of the superconducting gap decouples the motion of the condensate from microscopic degrees of freedom, and allows current in a superconductor to flow with vanishingly small dissipation. In the context of an LC oscillator, this means that we may straightforwardly fabricate superconducting resonators with very low decay rates κ , or equivalently very high quality factors $Q = \omega_0/\kappa$. In the language of particle physics or quantum optics, one might say that the line-width of a superconducting resonator¹ can be made extremely

¹That being said, we can and do use high- Q ($\sim 10^3 - 10^4$) resonators from normal conductors such as copper. While the utility of such a resonator will ultimately be limited by resistive loss, they too can be modeled by a single quantum degree of freedom.

narrow compared to the level spacing, such that individual energy levels may easily be resolved. The “quantum harmonic oscillator” is really describing the rigid oscillations of the *entire* superconducting condensate, shielded from decaying into the many-body environment by the superconducting gap.

2.2.2 Josephson junctions

As with classical circuits, capacitors and inductors have finite utility on their own. Even to make a single qubit, we must combine them with some nonlinearity: if we didn’t, all we could make are harmonic oscillators with evenly spaced energy levels. Without some sort of anharmonicity, we could never address the transition between two energy eigenstates, say the ground state $|0\rangle$ and first excited state $|1\rangle$, without accidentally climbing the harmonic oscillator ladder ($|2\rangle, |3\rangle, \dots$). To build an analogy to systems we learned about in introductory quantum mechanics, instead of building a circuit that makes harmonic oscillators, we would like to build a circuit that makes an “artificial atom” [33]: a non-linear quantum system with anharmonicity. If the anharmonicity in the energy level spacing is large enough, we can constrain ourselves to two levels to make an effective qubit.

When working with superconducting circuits, we realize this required anharmonic oscillator by replacing geometric inductors with Josephson junctions. A Josephson junction contains a thin insulating layer that splits the superconducting circuit but leaves them weakly coupled (Fig. (2.2a)). The split in the superconductor also splits the complex order parameter, and the order parameter on either side of the junction may have a different phase. However, the insulating barrier is thin enough that Cooper pairs may tunnel across the junction, weakly connecting the two sides.

Take the idealized picture shown in Fig. (2.2a): we have two (identical) superconducting

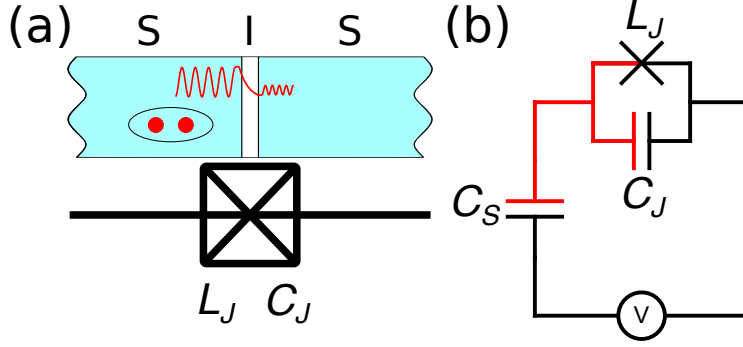


Figure 2.2: **The Josephson junction and the Cooper pair box:** (a) Schematic of a Josephson junction: two superconductors (S, light blue) are separated by a thin insulating barrier (I, white), across which Cooper pairs may tunnel. Josephson junctions act as nonlinear inductors (inductance L_J) and are typically represented by either a cross, or a cross in a box which signifies the inductance L_J in parallel with some small but finite capacitance C_J associated with the barrier. (b) The “Cooper pair box” circuit. Here, we’ve explicitly separated out the Josephson inductance and capacitance, and have added a shunt capacitance and a “voltage source.” Note that the red part of the circuit is only connected to the rest of the circuit via the junction, and can only exchange integer numbers of Cooper pairs with the rest of the circuit via Josephson tunneling.

islands separated by an insulating barrier forming a Josephson junction. If the temperature is much lower than the superconducting gap $k_B T \ll 2\Delta_{\text{BCS}}$, all single particle excitations are exponentially suppressed and the only free parameters in our system are n_L and n_R , the number of Cooper pairs on the left/right islands respectively. Since charge is conserved, we may simply consider the difference in Cooper pair number $n = n_L - n_R$, and label the state of the system as $|n\rangle$. We define the Cooper pair number operator:

$$\mathbf{n} = \sum_{n=-\infty}^{n=\infty} n |n\rangle \langle n| \quad (2.12)$$

which acts on a Cooper-pair number eigenstate $\mathbf{n} |n'\rangle = n' |n'\rangle$. Let’s build a phenomenological Hamiltonian that captures the physics of our idealized system: a Cooper pair can tunnel from left to right (taking $|n\rangle$ to $|n-1\rangle$), or from right to left (taking $|n-1\rangle$ to $|n\rangle$). Say that there’s some energy E_J associated with tunneling across the junction. A Hamiltonian

that describes the coherent tunneling of Cooper pairs can be written down as:

$$\mathbf{H}_J = -\frac{E_J}{2} \sum_{n=-\infty}^{n=\infty} |n-1\rangle \langle n| + |n\rangle \langle n-1| \quad (2.13)$$

The first term describes a Cooper pair tunneling from left to right ($n = n_L - n_R$ decreases by 1) and the second a Cooper pair tunneling from right to left. The sum in the Hamiltonian must run from $-\infty$ to ∞ , since we may have a disequilibrium of Cooper pairs biased to the right side of the junction just as easily as to the left side. More abstractly, this Hamiltonian describes (discrete) translations in the number n , and therefore we should expect the basis that diagonalizes this Hamiltonian to be the (discrete) Fourier conjugate to $|n\rangle$. Indeed, if we plug in the ansatz wave function

$$|\varphi\rangle = \sum_{n=-\infty}^{n=\infty} e^{in\varphi} |n\rangle \quad (2.14)$$

we find

$$\begin{aligned} \mathbf{H}_J |\varphi\rangle &= -\frac{E_J}{2} \sum_{n,n'} \left(|n-1\rangle \langle n| + |n\rangle \langle n-1| \right) e^{in'\varphi} |n'\rangle \\ &= -\frac{E_J}{2} \left(\sum_n e^{i(n+1)\varphi} |n\rangle + \sum_n e^{i(n-1)\varphi} |n\rangle \right) \\ &= -E_J \cos \varphi \sum_n e^{in\varphi} |n\rangle \\ \mathbf{H}_J |\varphi\rangle &= -E_J \cos \varphi |\varphi\rangle \end{aligned} \quad (2.15)$$

Since we said φ looks a lot like the Fourier conjugate to n , we could guess that we could write down $|n\rangle$ in the $|\varphi\rangle$ basis by taking an inverse Fourier transform of φ . Since n is discrete, φ

is only defined mod 2π :

$$\int_0^{2\pi} d\varphi e^{-in\varphi} |\varphi\rangle = \int_0^{2\pi} d\varphi \sum_{n'=-\infty}^{n'=\infty} e^{i(n'-n)\varphi} |n'\rangle = 2\pi \sum_{n'=-\infty}^{n'=\infty} \delta_{n,n'} |n'\rangle$$

$$|n\rangle = \frac{1}{2\pi} \int_0^{2\pi} d\varphi e^{-in\varphi} |\varphi\rangle \quad (2.16)$$

So we see, in fact, that our guess was correct. Given that φ is the Fourier conjugate of the Cooper pair number n , they form a set of canonically conjugate variables. Since $n = Q/2e$, where e is the electron charge, we see that φ is directly proportional to the node flux ϕ across the junction that we defined earlier

$$\varphi = \frac{\phi}{\Phi_0} \bmod 2\pi$$

$$\Phi_0 = \frac{\hbar}{2e} \quad (2.17)$$

where Φ_0 is the reduced superconducting flux quantum. Reminding ourselves of the definition of ϕ in terms of the voltage between two nodes, this equivalence leads to the second Josephson relationship Eqn. 2.1(b). The equivalence between the node flux ϕ and the phase difference φ leads to a physical interpretation of Josephson junctions that will guide our discussion: Eqn. 2.15 implies that a Josephson junction acts like a non-linear inductor whose energy, rather than being quadratic in the flux φ across the element, is $\propto \cos \varphi$.

2.2.3 The Cooper pair box

The above picture of a Josephson junction was oversimplified: we failed to take into account the electrostatic energy associated with a charge imbalance between the two sides of the

Josephson junction. Examining Figure 2.2(a), we see that in addition to being a nonlinear inductor, a Josephson junction consists of two metal electrodes that will have some mutual capacitance C_J . Given our prior work quantizing the LC oscillator, it will behoove us to consider electrostatic energy of Cooper-pair disequilibrium in terms of the charging of the capacitor formed by the two superconducting leads.

Consider the setup described in Figure 2.2(b): here, we've added in a shunt capacitor C_S which completely isolates the charge carriers in the red conductor (the “Cooper pair box”) from the charge carriers in the rest of the circuit except by tunneling through the junction. We also have a voltage source in the black part of the circuit, which serves to bias the red island relative to the black part of the circuit via the shunting capacitor. This voltage source may either describe a deliberately applied electric field or the microscopic electrostatic environment of the junction.

In this circuit, we see that any tunneling events will cause a buildup of charge on the capacitor plates. If we write $C_\Sigma = C_S + C_J$, we can write down the electrostatic contribution to the Hamiltonian that comes from the capacitance C_Σ

$$\mathbf{U} = 4E_c(\mathbf{n} - n_g)^2 \tag{2.18}$$

$$E_c = \frac{e^2}{2C_\Sigma}$$

$$n_g = -\frac{C_g V}{2e}$$

Here, E_C is the charging energy of one electron, i.e. the electrostatic energy associated with taking an electron from one side of the capacitor and moving it to the other (the factor of

4 comes from the fact that we're dealing with Cooper pairs that have $q = 2e$.) We have encapsulated the voltage source within an *offset charge* n_g : a change in V (either deliberately or by microscopic fluctuations) will shift the equilibrium charge distribution, and thus modify the Hamiltonian. Note that while the number of Cooper pairs on the island is quantized (by virtue of our capacitor plates being isolated except through the Junction, which only integer numbers of Cooper pairs can tunnel across), n_g is a continuous variable.

Taking into account both the electrostatic energy and the Josephson tunneling energy, we write down the full Hamiltonian for our first prototype qubit: the Cooper pair box (CPB)

$$\mathbf{H}_{CPB} = 4E_c(\mathbf{n} - n_g)^2 - E_J \cos \varphi \quad (2.19)$$

If we write the number operator in the phase basis, $\mathbf{n} = i\partial/\partial\varphi$, the time independent Schrödinger equation for this Hamiltonian becomes a differential equation called Mathieu's equation, and is exactly solvable in terms of Mathieu functions [34]. Plotted in Figure 2.3(a) are several low-lying eigenvalues of the CPB Hamiltonian as a function of the offset charge n_g for the ratio of Hamiltonian parameters $E_J/E_C = 1$.

The CPB was, in fact, the first superconducting quantum bit in the sense that it was the first superconducting circuit in which it was demonstrated that a coherent superposition of two macroscopic states of the system could be prepared [29]. The nonlinearity of the Josephson junction shows up clearly in Figure 2.3(a): the energy levels are far from equally spaced at all values of n_g . Examining Figure 2.3(a), one sees the remnants of the parabolic dependence of the energy bands on offset charge for each integer value of Cooper pair. In the absence of tunneling, each parabola would be completely uncoupled to the adjacent parabola, however the CPB spectrum has a strong avoided crossing at $n_g = 1/2$ caused

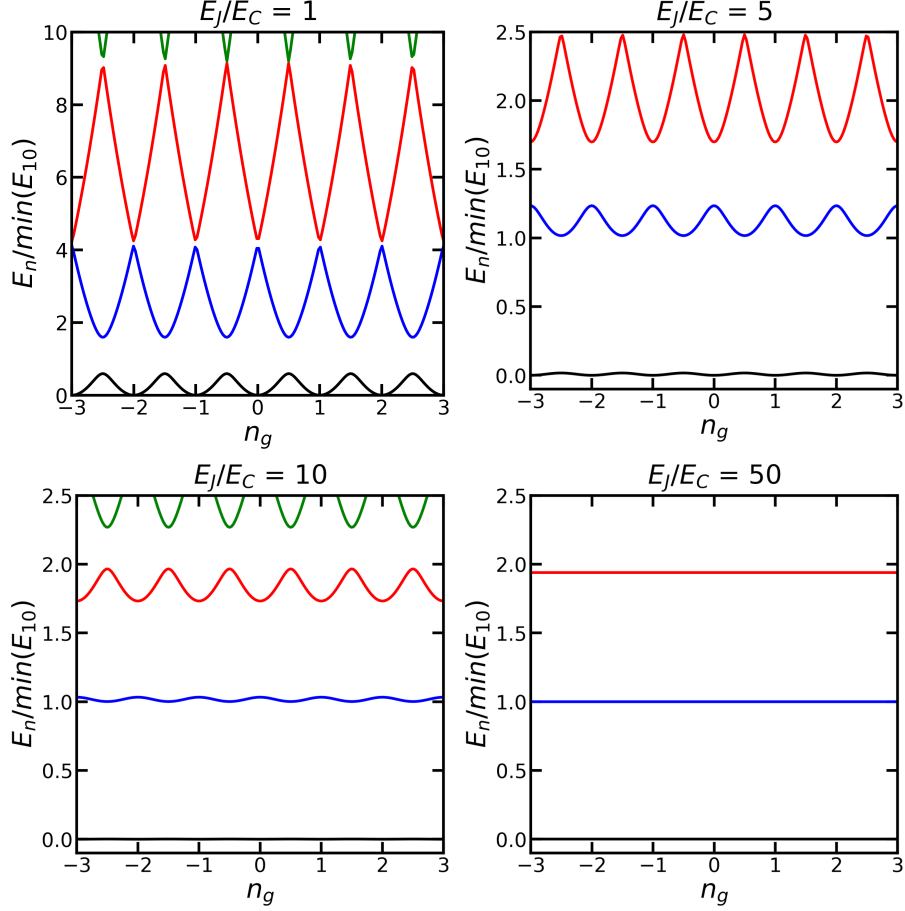


Figure 2.3: **Eigenvalues of the Cooper pair box Hamiltonian:** Eigenvalues of the Cooper pair box Hamiltonian, as a function offset charge n_g for various ratios E_J/E_C . As the ratio E_J/E_C goes up, charge dispersion is exponentially suppressed, but anharmonicity is also suppressed. Reproduced from Ref. [34].

by the Josephson tunneling term. In Ref. [29], coherent oscillations between Cooper-pair number eigenstates were observed by nonadiabatically tuning the gate voltage from a point far from $n_g = 1/2$, where the CPB was likely to be in the $|n\rangle = |0\rangle$ state, to $n_g = 1/2$ where the eigenstates of the CPB Hamiltonian are $(|0\rangle + |1\rangle)/\sqrt{2}$. This was a proof of principle that coherent manipulations could be executed on the two lowest energy levels of the CPB. This, coupled with proposals to read out the charge eigenstates of the CPB using a single electron transistor [35], was at the time a promising route of pursuing solid state quantum computing.

2.2.4 Energy scales

Up until this point, we have written down energies E_J and E_C without considering the scale of these energies. To progress further, we will need to consider the characteristic energy scale of Josephson junction circuits. This will also define a characteristic temperature at which superconducting circuit experiments must be performed, and prelude some of the experimental considerations we will encounter in Chapter 3.

The only energy scale that has entered the picture so far is the superconducting gap² $\Delta_{\text{BCS}}(T = 0) \approx 1.76k_bT_c$, where T_c is the critical temperature of the superconductor [21]. Pragmatically, coherent control of a qubit system with $|0\rangle \rightarrow |1\rangle$ transition energy $(E_1 - E_0) = \hbar\omega_{01}$ is achieved by applying coherent electromagnetic radiation at frequency ω_{01} to the circuit. This presents a hard upper bound to the characteristic energies of the system $\hbar\omega_{01} \ll 2\Delta_{\text{BCS}}$, else the control signals at ω_{01} will carry enough energy per-photon to break Cooper pairs and cause dissipation [21]. Fortunately, it is fairly straightforward to engineer Josephson junctions such that $E_J \leq \Delta_{\text{BCS}}$.

For commonly used aluminum, $\Delta_{\text{BCS}}/h \approx 40$ GHz, and thus superconducting circuits made from aluminum should be operated at frequency $\ll 80$ GHz. Modern telecommunication technology has made microwave electronics in the 3 – 10 GHz range readily available at reasonable expense, and engineering microwave setups at these frequencies is somewhat more forgiving than at higher frequencies. Thus, we will quote the “typical” energy of a superconducting qubit $\omega_{01}/(2\pi) \sim 5$ GHz.

In addition to anharmonicity, another prerequisite for coherently operating a quantum circuit is that the system of interest can be reliably initialized in a known fiducial state,

²The temperatures we work at in this thesis are all much less than the superconducting transition temperature T_c of aluminum, where $\Delta_{\text{BCS}}(T) \approx \Delta_{\text{BCS}}(0)$. In light of this, I’ll drop the T dependence of Δ_{BCS} , and quote only the zero temperature value.

and that spurious excitations over the course of an experiment are rare. Both of these prerequisites are satisfied if the temperature of the experiment is low enough such that $k_b T \ll \hbar \omega_{01}$. In this case, thermal excitations of the qubit are exponentially suppressed in temperature, and initializing the qubit into the $|0\rangle$ simply becomes a matter of waiting for it to thermalize with its environment. Converting to temperature, a transition frequency of 5 GHz corresponds to $T \approx 240$ mK: superconducting qubit experiments require rather cold temperatures! Fortunately, commercially available dilution refrigerators routinely reach base temperatures < 10 mK, and thus an ambient temperature that satisfies $k_b T \ll \hbar \omega_{01}$ is fairly easy to achieve.³

2.2.5 The transmon

A major problem with early implementations of the CPB is apparent from Figure 2.3: for low E_J/E_C the qubit frequency ω_{01} is a strong function of the offset charge n_g . The information in a qubit is not only encoded in the relative amplitudes of $|0\rangle$ and $|1\rangle$, but also in the relative *phase*. Recall how a quantum two-level system initially in $|\psi(0)\rangle = (|0\rangle + |1\rangle)/\sqrt{2}$ evolves as a function of time

$$|\psi(t)\rangle = \frac{1}{\sqrt{2}}(|0\rangle + e^{-i\omega_{01}t} |1\rangle) \quad (2.20)$$

If ω_{01} is constant, this phase factor evolves deterministically, and we may account for it by working in the frame rotating with the qubit state vector. However, we see that, for low E_J/E_C , ω_{01} is *not* constant: it depends strongly on the offset charge n_g , and changes by large values over a fraction of a Cooper pair. As stated previously, the offset voltage V also

³As we shall see in future chapters, cooling everything down to the nominal temperature of the cryostat isn't so straightforward.

serves to model the local electrostatic environment of the CPB. It is clear then that small fluctuations in the electrostatic environment change the transition frequency ω_{01} , and thus cause the relative phase of $|0\rangle$ and $|1\rangle$ to evolve nondeterministically. This implies that CPB qubits are extremely sensitive to charge noise.

From Fig. 2.3, it is abundantly clear that a solution to this problem is to work in the large E_J/E_C regime, where the charge dispersion is suppressed. We can accomplish this by making the shunt capacitance C_S very large compared to the intrinsic Josephson junction capacitance C_J , driving the capacitive charging energy E_C down and the ratio E_J/E_C up. A CPB in the parameter range $E_J/E_C \approx 50 - 70$ is called a “transmon” qubit [34], and will be the main superconducting circuit featured in this thesis.

One immediate drawback we can see from Fig. 2.3 is that in the large E_J/E_C regime the anharmonicity becomes weaker: as we increase E_J/E_C , the circuit begins to look more and more like a harmonic oscillator. To understand why transmons have low anharmonicity, we can map the Hamiltonian Eqn. (2.19) onto a more familiar problem: a rigid pendulum of length l and mass m experiencing a gravitational force g [34, 31]. The Hamiltonian for this system is

$$\mathbf{H}_{pend} = \frac{\mathbf{L}_z^2}{2ml^2} - mlg \cos \varphi \quad (2.21)$$

Here, φ becomes the angle that the pendulum makes with respect to its equilibrium position (consistent with our demand that φ only be defined mod 2π), and (ignoring the offset charge for now) $\mathbf{n} = \mathbf{L}_z/\hbar$ becomes the conjugate (angular) momentum.

In this model, $E_C = \hbar^2/8ml^2$ is inversely proportional to the pendulum’s moment of inertia, and $E_J = mgl$ is the magnitude of the gravitational potential energy. In the transmon

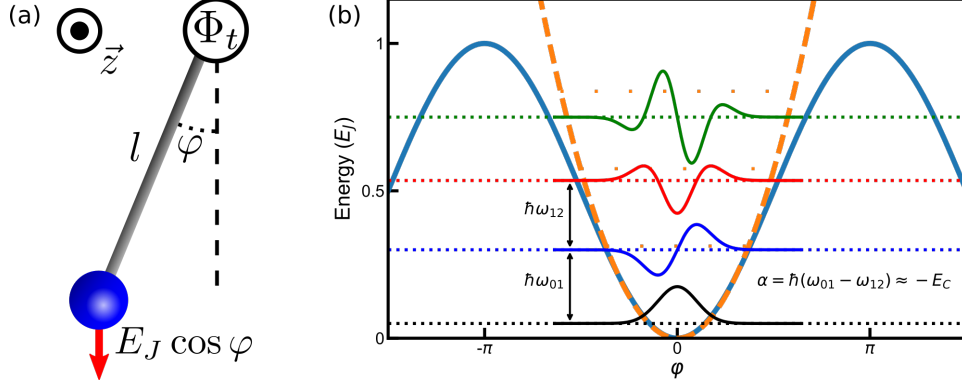


Figure 2.4: **A mechanical analogue to the Cooper pair box Hamiltonian:** (a) An analogue to the Cooper pair box Hamiltonian: a rigid pendulum of length l and mass m living in a gravitational field. To account for the offset charge n_g , we also give the pendulum some fictitious charge q and tack on a vector potential to the Hamiltonian by threading a flux Φ_t through a tube at the pendulum's hinge. (b) Josephson-junction cosine potential (light blue) and eigenstates (eigenenergies dotted, wave functions solid) compared to a harmonic oscillator potential (dashed orange) and eigenenergies (dotted orange). In the transmon regime, large E_J makes the circuit act like an “almost harmonic oscillator” (i.e. a quantum pendulum in the small angle regime) and exponentially suppresses tunneling to into neighboring wells (a complete loop of the pendulum) through which charge noise enters. Inspired by Figure 3 of [34].

regime $E_J/E_C \gg 1$, the mechanical Hamiltonian is dominated by the gravitational potential energy term, and we expect the rotor to remain in the almost harmonic small angle regime.

In the small angle limit, we may expand the cosine in the CPB Hamiltonian

$$\mathbf{H} = 4E_C \mathbf{n}^2 - E_J \left(1 - \frac{\varphi^2}{2} + \mathcal{O}(\varphi^4)\right) \approx 4E_C \mathbf{n}^2 + \frac{E_J \varphi^2}{2} \quad (2.22)$$

where we've dropped both the constant offset and all of the anharmonic higher order terms.

This is exactly the quantum harmonic oscillator we considered in §2.2.1, with $\mathbf{n} = \mathbf{Q}/(2e)$

and $\varphi = \phi/\Phi_0$, and we may recycle the machinery from that discussion:

$$\mathbf{H} \approx \frac{\mathbf{Q}^2}{2C_\Sigma} + \frac{\phi^2}{2L_J} \quad (2.23a)$$

$$L_J = \frac{\Phi_0^2}{E_J} \quad (2.23b)$$

Here, we reintroduce the variable C_Σ and introduce the Josephson inductance L_J . We then write down the Hamiltonian as $\hbar\omega_J(\mathbf{a}^\dagger\mathbf{a} + 1/2)$, with the raising and lowering operators defined exactly as in section 2.2.1. We may just copy and paste the results we already obtained:

$$\omega_J = \sqrt{\frac{1}{L_J C_\Sigma}} = \frac{1}{\hbar} \sqrt{8E_J E_C} \quad (2.24a)$$

$$Z_J = \sqrt{\frac{L_J}{C_\Sigma}} = \frac{\hbar}{e^2} \sqrt{\frac{E_C}{2E_J}} \quad (2.24b)$$

$$\phi_{ZPF} = \sqrt{\frac{\hbar Z_J}{2}} = \Phi_0 \left(\frac{2E_C}{E_J} \right)^{\frac{1}{4}} \quad (2.24c)$$

$$Q_{ZPF} = \sqrt{\frac{\hbar}{2Z_J}} = 2e \left(\frac{E_J}{32E_C} \right)^{\frac{1}{4}} \quad (2.24d)$$

Eqn. (2.24a) defines the so called “Josephson plasmon frequency”, which is simply the frequency of harmonic oscillation of a Josephson junction modeled as a parallel capacitance and inductance. Z_J is the characteristic impedance of the junction: for a sense of scale, $\hbar/e^2 \approx 4 \text{ k}\Omega$. Eqns. (2.24c) and (2.24d) are the magnitude of the zero point flux and charge fluctuations respectively, with the proper normalization factors to convert $\phi \leftrightarrow \varphi$ and $\mathbf{Q} \leftrightarrow \mathbf{n}$ pulled out. From these relations, we can immediately read off that, for $E_J/E_C \gg 1$ the zero point fluctuations in the charge are larger than a Cooper pair, and the zero point fluctuations in the phase are small. Since this entire argument hinges on the small angle

approximation, it is comforting that ϕ_{ZPF} is less than unity. It is interesting to note that, even though the transmon is derived from the CPB “charge” qubit⁴, the Cooper pair number n clearly isn’t a good quantum number. The physical qubit states are the plasma oscillations of the superconducting condensate across the Josephson junction.

The small anharmonicity is, however, crucial to our ability to use transmons as a qubit. To add the anharmonicity back into the equation, consider the $\mathcal{O}(\varphi^4)$ term as a perturbation to the harmonic oscillator Hamiltonian:

$$\mathbf{H}^{(4)} = -\frac{E_J \varphi^4}{4!} = -E_J \frac{1}{24} \left(\frac{\Phi_{ZPF}}{\Phi_0} (\mathbf{a}^\dagger + \mathbf{a}) \right)^4 \quad (2.25)$$

where I’ve used the rules laid out in Section 2.2.1 to write down φ in terms of the raising and lowering operators. After some algebra, we reduce this to terms diagonal in the unperturbed Hamiltonian⁵:

$$\mathbf{H}^{(4)} = -\frac{E_C}{2} (\mathbf{a}^\dagger \mathbf{a}^\dagger \mathbf{a} \mathbf{a} + 2 \mathbf{a}^\dagger \mathbf{a}) \quad (2.26)$$

Then the total Hamiltonian, to the level of first order perturbation theory, is

$$\mathbf{H} \approx (\hbar\omega_J - E_C)(\mathbf{a}^\dagger \mathbf{a} + 1/2) - \frac{E_C}{2} \mathbf{a}^\dagger \mathbf{a}^\dagger \mathbf{a} \mathbf{a} \quad (2.27)$$

The first excited state energy is renormalized to

$$\hbar\omega_{01} \approx \sqrt{8E_J E_C} - E_C \quad (2.28)$$

⁴Superconducting qubits tend to be classified by their good quantum numbers, i.e. what basis the Hamiltonian is diagonalized in during normal operation. The CPB is a charge qubit in that it is intended to be operated in a parameter regime where the Hamiltonian is diagonalized in the charge basis (away from the avoided crossings). Phase and flux qubits are also common classifications.

⁵Terms non-diagonal in the unperturbed Hamiltonian may be ignored because of the rotating wave approximation, see Appendix A. This is only valid if $E_C \ll \hbar\omega_J$, which is always true in the transmon regime.

and the second term gives us the anharmonicity $\alpha = \omega_{12} - \omega_{01} \approx -E_C$. As stated, before, the desired $E_J/E_C \sim 50$ ratio is yielded by increasing the shunt capacitance C_S to drive down E_C . Typically $C_S \sim 70$ fF, while the intrinsic Josephson capacitance is typically $\lesssim 1$ fF. Fortunately, for $C_\Sigma \approx 70$ fF, $E_C/h \approx 300$ MHz, and thus for a transmon qubit with $\omega_{01}/(2\pi) \sim 5$ GHz, the anharmonicity remains an appreciable fraction of the transition frequency, and sufficiently large to address only the two-level subspace $|0\rangle$ and $|1\rangle$.

2.2.6 Where did the charge noise go?

To build intuition for why transmon qubits are insensitive to charge noise, let's go back to the quantum pendulum analogue. Remember that, when we made the analogy with the pendulum, we ignored the offset charge n_g in the CPB Hamiltonian. To add the charge offset into the mechanical analogue, we may give the rotor a fictitious charge q and thread a magnetic flux Φ_t through a tube at the hinge of the pendulum. To be clear, the magnetic field in this toy model only exists inside the tube: the rotor (with its fictitious charge positioned at the end of the pendulum) never experiences any Lorentz forces. However, the presence of the magnetic flux will cause a nonzero vector potential in the vicinity of the charge⁶:

$$\vec{A} = \frac{\Phi_t}{2\pi r} \vec{z} \times \vec{r} \quad (2.29)$$

The presence of the vector potential takes $\vec{p} \rightarrow \vec{p} - q\vec{A}$, which takes the angular momentum to

$$\vec{L}_z = (\vec{r} \times \vec{p})_z - q(\vec{r} \times \vec{A})_z \Rightarrow \hbar \left(-i \frac{\partial}{\partial \phi} - \frac{q}{4\pi e} \frac{\Phi_t}{\Phi_0} \right) \quad (2.30)$$

⁶It's worth repeating this for clarity, since things can get confusing here: the charge q on the pendulum is a completely fictitious charge living within the picture of the rotor analogue, with no relation to the actual charges in the superconductor. It exists solely to interact with the fictitious flux Φ_t

If we plug the new angular momentum into Eqn. (2.21) and compare it to the original Cooper pair box Hamiltonian, we immediately see that

$$n_g \propto \Phi_t \tag{2.31}$$

This is curious. Classically, the magnetic flux completely confined to the tube at the hinge of the rotor could never interact with the charge at the end of the pendulum. However, quantum mechanically, we know that a charged particle that traverses around a magnetic flux may pick up an Aharonov-Bohm phase, even if it never experiences a Lorentz force [36]. Therefore, the only way for the particle to possibly be influenced by the flux Φ_t is for it to completely traverse the loop. Since a simple phase doesn't effect the spectrum of a quantum system, the only way the qubit spectrum can depend on Φ_t is via *interference between the path where the rotor didn't do a loop and the path where it did*.

In the transmon regime $E_J \gg E_C$, however, we have seen that at low excitation number our rotor will be well confined to small angles: the “force of gravity” is strong and the “mass” of the rotor (moment of inertia) is large. In order for the rotor to undergo a so called phase slip (a loop around the flux that changes the equilibrium value of ϕ by 2π), the particle (with mass $\propto 1/E_C$) must tunnel over a barrier whose height is proportional to E_J . However, the probability of tunneling over a large potential energy barrier is *exponentially* suppressed in the barrier height ($E_J \gg E_C$), and since the only way the energy eigenvalues can depend on n_g is via interference between paths with different winding numbers around the hinge, the transmon's dependence on the offset charge n_g is also exponentially suppressed [34].

Transmon qubits were first proposed in Ref. [34] as a charge qubit capacitively coupled to a coplanar waveguide resonator, and were quickly used to experimentally demonstrate (at the

time) impressive results in the dispersive regime of circuit quantum electrodynamics [37, 38, 39]. Since then, material and design parameters have improved to the point where transmon qubits commonly have coherence times approaching or exceeding $100\ \mu s$ [40, 41, 42, 43, 44]. As such, the transmon qubit has become a workhorse platform for demonstrating proof of principal quantum computing applications [45, 46, 47, 48, 49, 50, 51, 52, 53, 54] and for conducting “microwave quantum optics” experiments [55, 56, 57, 58, 59, 43]. While transmon qubits have rapidly advanced in coherence to the point where many complex protocols can be performed on them, it is clear that further suppression of loss mechanisms is needed for scalable quantum computing applications, and research on decoherence mechanisms in superconducting qubits remains an active field. We will discuss decoherence in greater detail in Chapter 4.

2.3 Circuit quantum electrodynamics

While making the transmon insensitive to charge noise greatly increases possible coherence times, it also invalidates early proposals for control and readout mechanisms of the CPB, which relied heavily on the different electrostatic configurations of various Cooper pair number states. The transmon qubit is, however, exponentially insensitive to its local electrostatic environment, rendering these schemes useless. Therefore, we need a new scheme to read out and manipulate transmon qubits.

Circuit quantum electrodynamics (cQED) [31, 33, 60, 61, 62, 63, 64], is an experimental protocol which allows us to manipulate and read out charge insensitive qubits while simultaneously protecting the qubit from what would be its dominant decay channel: spontaneous decay into the electromagnetic vacuum. In cQED, we place the qubit into a high finesse

microwave cavity with resonant frequency $\omega_c \sim \omega_{01}$. This has two advantages: the cavity shields the electromagnetic environment of the qubit, greatly reducing the density of states available for the qubit to decay into [65, 40], while the confined discrete electromagnetic mode of the cavity strongly couples to the qubit, allowing for fast manipulation and readout.

2.3.1 Putting together the building blocks of cQED

In essence, cQED is the combination of two of the basic components we’ve already built up: a harmonic oscillator (which we’ll model as an LC circuit) and a qubit (which we’ll take to be a transmon.) The model here roughly follows [66, 64, 67]: in real experiments, the harmonic oscillator is replaced with an electromagnetic cavity, which has infinite discrete modes (see Chapter 3 for more details.) However, if the detuning Δ between the qubit and the fundamental mode of the cavity ($|\Delta| = |\omega_{01} - \omega_c|$, where ω_c is the cavity frequency) is much smaller than the detuning between the qubit and any other mode, the single mode model considered here is an excellent approximation.

2.3.1.1 Driving the system

We begin by considering the circuit in Fig. 2.5 composed of a transmon qubit (blue) capacitively coupled to both a driving field (black) and a harmonic oscillator (orange.) To start off, let’s ignore the orange harmonic oscillator and see how the driving field couples to the qubit. In this sub-circuit, we identify two nodes with node fluxes ϕ_1 and ϕ_D . We can write down the Lagrangian for this subcircuit with ease by adding up the capacitive and inductive energies component-wise

$$\mathcal{L} \approx \frac{1}{2}C_\Sigma \dot{\phi}_1^2 + \frac{1}{2}C_D(\dot{\phi}_1 - \dot{\phi}_D)^2 - \frac{1}{2L_J}\phi_1^2 \quad (2.32)$$

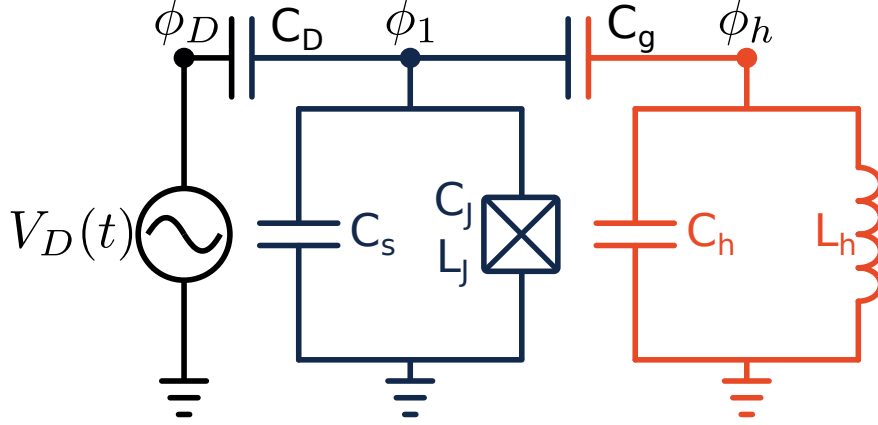


Figure 2.5: **Schematic of circuit QED:** A circuit schematic of the circuit quantum electrodynamics (cQED) protocol. A qubit (blue) is coupled to a microwave frequency electromagnetic cavity, modeled as an LC oscillator (orange). We apply microwave qubit control and readout pulses through an external drive (modeled as a voltage source in black.) Here, we take the qubit to be a transmon qubit, and the coupling to the cavity/drive to be capacitive.

where, for now, we've dropped the (small) nonlinearity of the Josephson junction potential. The Lagrangian at first appears to be a function of two dynamic variables, however we notice immediately that $\dot{\phi}_D = V(t)$, and thus ϕ_D is completely constrained by the applied input voltage. To find the canonical momentum to ϕ_1 , let's set $V_D(t) = 0$ and write down the Hamiltonian

$$H \approx \frac{Q_1^2}{2(C_\Sigma + C_D)} + \frac{\phi_1^2}{2L_J} \quad , \quad Q_1 = \dot{\phi}_1(C_\Sigma + C_D) \quad (2.33)$$

The canonical charge (and the capacitive energy) is renormalized by the presence of the coupling capacitance, as we should expect. Now, the terms in the Lagrangian that incorporate the presence of the drive are

$$\mathcal{L}_d = \frac{1}{2}C_D\dot{\phi}_D^2 - C_D\dot{\phi}_D\dot{\phi}_1 = \frac{1}{2}C_DV_D(t)^2 - C_DV_D(t)\dot{\phi}_1 \quad (2.34)$$

The first term contains no dynamic variables, and we may ignore it. The second term

encapsulates the effect of the drive on the qubit: if we assume weak coupling to the drive, we may write it down in terms of the canonical charge and add it back into the Hamiltonian as a perturbation

$$H = \frac{Q_1^2}{2(C_\Sigma + C_D)} + \frac{\phi_1^2}{2L_J} - \frac{C_D}{C_\Sigma + C_D} V_D(t) Q_1 \quad (2.35)$$

As we have done multiple times now, we quantize this Hamiltonian in the harmonic oscillator basis, taking $\mathbf{Q}_1 = -iQ_{ZPF}(\mathbf{a}_q - \mathbf{a}_q^\dagger)$, where $\mathbf{a}_q^\dagger$ ($\mathbf{a}_q$) is the raising (lowering) operator for the transmon circuit. Doing so gives us

$$\mathbf{H} = \hbar\omega_q(\mathbf{a}_q^\dagger\mathbf{a}_q + \frac{1}{2}) + \epsilon(t)(\mathbf{a}_q - \mathbf{a}_q^\dagger) \quad (2.36)$$

where $\omega_q = (L_J(C_\Sigma + C_D))^{-1/2} \approx (L_J C_\Sigma)^{-1/2}$ for weak coupling to the driving and

$$\epsilon(t) = \frac{iQ_{ZPF}C_D V_D(t)}{C_\Sigma + C_D} \quad (2.37)$$

is the amplitude of the drive. This analysis gives us a mathematical framework by which we may analyze how an applied (classical) electric field⁷ couples to a qubit (or a harmonic oscillator.) In fact, this framework of driving the qubit with classical electric fields is the basis of coherent control, where we use classical fields to manipulate the state of a quantum system. We will discuss coherent control more in Chapter 3.

2.3.1.2 Coupling a qubit and a harmonic oscillator

We now turn to the case of the qubit coupled to the harmonic oscillator. The analysis here is exactly the same, however instead of a classical drive, we note that even when the harmonic

⁷In the literature, the applied field term is typically written down as $\propto \mathbf{a}_q^\dagger + \mathbf{a}_q$. The sign change is simply a convention that depends on whether we took Q or ϕ to be the canonical coordinate.

oscillator is in its ground state, there are zero point voltage fluctuations. To see this, we can simply take the time derivative of the flux operator ϕ_h using the Heisenberg equations of motion [31]

$$\mathbf{V}_h(t) = \frac{d\phi_h}{dt} = \frac{i}{\hbar}[\mathbf{H}, \phi_h] = \frac{\mathbf{Q}_h}{C_h} = -iV_{ZPF}(\mathbf{a}_c - \mathbf{a}_c^\dagger) \quad (2.38)$$

$$V_{ZPF} = \frac{Q_{ZPF}}{C_h} = \omega_c \sqrt{\frac{\hbar Z_0}{2}} \quad (2.39)$$

where $\mathbf{a}_c^\dagger$ ($\mathbf{a}_c$) are the raising (lowering) operators for the harmonic oscillator cavity mode, and $\omega_c = (L_h C_h)^{-1/2}$ is the harmonic oscillator resonant frequency. These zero point voltage fluctuations will act on the qubit in the exact same way as the drive voltage $V_D(t)$. Thus, even when the Harmonic oscillator is in the ground state, there will be some interaction between the qubit and the harmonic oscillator. Adding back in the Josephson non-linearity, we may write down the composite transmon + harmonic oscillator + coupling Hamiltonian in the same form as Eqn. 3.12, replacing the applied voltage with the zero-point voltage fluctuations of the cavity

$$\begin{aligned} \mathbf{H} &= \left[\hbar\omega_q(\mathbf{a}_q^\dagger \mathbf{a}_q + \frac{1}{2}) + \mathcal{O}(\phi_1^4) \right] + \hbar\omega_c(\mathbf{a}_c^\dagger \mathbf{a}_c + \frac{1}{2}) - \frac{C_g}{C_\Sigma + C_g} \mathbf{V}_h(t) \mathbf{Q}_1(t) \\ &= \left[\hbar\omega_q(\mathbf{a}_q^\dagger \mathbf{a}_q + \frac{1}{2}) + \mathcal{O}(\phi_1^4) \right] + \hbar\omega_c(\mathbf{a}_c^\dagger \mathbf{a}_c + \frac{1}{2}) - \hbar g(\mathbf{a}_q - \mathbf{a}_q^\dagger)(\mathbf{a}_c - \mathbf{a}_c^\dagger) \end{aligned} \quad (2.40)$$

where I have absorbed all the capacitances and zero-point fluctuations into the constant g known as the vacuum Rabi constant. We now make two approximations:

1. We invoke the nonlinearity of the transmon circuit to constrain ourselves to the qubit

manifold, i.e the transmon $|0\rangle$ and $|1\rangle$ states. When we think of quantum two-level systems, the canonical example that comes to mind is a spin- $1/2$ system. We therefore import the language of spins to model the effective two-level transmon qubit, and replace the harmonic oscillator-like operators with Pauli operators: the qubit raising operator $\mathbf{a}_q^\dagger$ becomes $\boldsymbol{\sigma}_+$, the lowering operator $\mathbf{a}_q$ becomes $\boldsymbol{\sigma}_-$, and the number operator $\mathbf{a}_q^\dagger \mathbf{a}_q + 1/2$ maps onto $\boldsymbol{\sigma}_z/2$.

2. We invoke the rotating wave approximation and say that, if $g \ll \omega_q, \omega_c$ counter-rotating terms such as $\mathbf{a}_c^\dagger \mathbf{a}_q^\dagger$ (or, equivalently, $\mathbf{a}_c^\dagger \boldsymbol{\sigma}_+$) will oscillate quickly over the timescale of system evolution, and will thus average out to zero (see Appendix A for a more complete discussion of the rotating wave approximation). We can therefore drop the $\mathbf{a}_c^\dagger \boldsymbol{\sigma}_+$ and $\mathbf{a}_c \boldsymbol{\sigma}_-$ terms.

Upon making these approximations, we arrive at the famous Jaynes-Cummings Hamiltonian, describing the interaction of a two-level atom with an electromagnetic mode:

$$\mathbf{H}_{\text{JC}} = \frac{\hbar\omega_q}{2} \boldsymbol{\sigma}_z + \hbar\omega_c(\mathbf{a}_c^\dagger \mathbf{a}_c + \frac{1}{2}) + \hbar g(\mathbf{a}_c^\dagger \boldsymbol{\sigma}_- + \mathbf{a}_c \boldsymbol{\sigma}_+) \quad (2.41)$$

2.3.2 The Jaynes-Cummings model

The Jaynes-Cummings (JC) model is a powerful theoretical tool for exploring the interaction of electromagnetic radiation with superconducting qubits. When deriving it, we already got a taste of how we may manipulate qubits with coherent electromagnetic drives (see Chapter 3 for a more complete discussion of coherent control.) We will now see how, within the framework of the JC model, we may also measure the state of the qubit.

The Jaynes-Cummings Hamiltonian consists of three terms: the bare qubit energy, the

bare cavity (harmonic oscillator) energy, and the interaction term. We may think of the interaction term as the action of the qubit exchanging an excitation with the cavity ($\mathbf{a}_c^\dagger \boldsymbol{\sigma}_-$) and it's Hermitian conjugate, the cavity exchanging an excitation with the qubit ($\mathbf{a}_c \boldsymbol{\sigma}_+$). The interaction term conserves the *total* number of quanta in the system, so total excitation number $\mathbf{N} = -\boldsymbol{\sigma}_z/2 + \mathbf{a}^\dagger \mathbf{a}$ is a constant of the motion and we may subtract it off. Doing so yields

$$\mathbf{H}' = \mathbf{H}_{\text{JC}} - \hbar\omega_c \mathbf{N} = \frac{\hbar\Delta}{2} \boldsymbol{\sigma}_z + \hbar g (\mathbf{a}^\dagger \boldsymbol{\sigma}_- + \mathbf{a} \boldsymbol{\sigma}_+) \quad (2.42)$$

where $\Delta = \omega_q - \omega_c$ is the cavity/qubit detuning. Since excitation number is conserved, the Hamiltonian is block diagonal with 2×2 blocks connecting the states $|n, 0\rangle$ and $|n-1, 1\rangle$, where the quantum numbers are the resonator excitation number and the qubit state respectively. We can write down the block of the Hamiltonian that acts on states in the space $\{|n-1, 1\rangle, |n, 0\rangle\}$ as a 2×2 matrix

$$\mathbf{H}'/\hbar = \begin{pmatrix} \Delta/2 & \sqrt{n}g \\ \sqrt{n}g & -\Delta/2 \end{pmatrix} \quad (2.43)$$

This Hamiltonian is straightforward to diagonalize, and yields two coupled qubit-photon eigenstates

$$|+, n\rangle = \cos \theta_n |n-1, 1\rangle + \sin \theta_n |n, 0\rangle \quad (2.44a)$$

$$|-, n\rangle = -\sin \theta_n |n-1, 1\rangle + \cos \theta_n |n, 0\rangle \quad (2.44b)$$

$$\theta_n = \frac{1}{2} \tan^{-1} \left(\frac{2g\sqrt{n}}{\Delta} \right) \quad (2.44c)$$

with energy levels

$$E_{\pm,n} = \pm \hbar \sqrt{ng^2 + (\Delta/2)^2} \quad (2.45)$$

There are several important features of this result to highlight:

1. There are only two parameters that characterize the system: the cavity/qubit detuning Δ , and the cavity/qubit coupling $g\sqrt{n}$, which depends on the number of quanta (photons) in the cavity.
2. When $\Delta = 0$ (i.e. the cavity and the qubit are on resonance), $\theta_n = 45^\circ$, and the eigenstates become equal and orthogonal superpositions of $|n-1, 1\rangle$ and $|n, 0\rangle$. These states are *not* degenerate: they undergo an avoided crossing and are split by $2\hbar g\sqrt{n}$, as shown in the Fig. 2.6. These states are neither qubit nor cavity states: they are hybridized states, consisting of both cavity and qubit in equal weight. When there is only one excitation in the system ($n = 1$), we can interpret this splitting as a quantum of energy being freely exchanged between the harmonic oscillator and the qubit at a rate $2g$, hence the name Vacuum Rabi splitting.
3. In the limit where the $\Delta \gg g\sqrt{n}$, the eigenstates of the Hamiltonian are *almost* pure qubit/cavity states, but not quite. We take this to mean that for example, a qubit excitation spends some time as a photon in the cavity. Along the same lines, the energy eigenvalues of the Hamiltonian are *almost* the bare qubit and cavity energies ω_q and ω_c , but not quite. We call these shifted eigenstates “dressed” states.

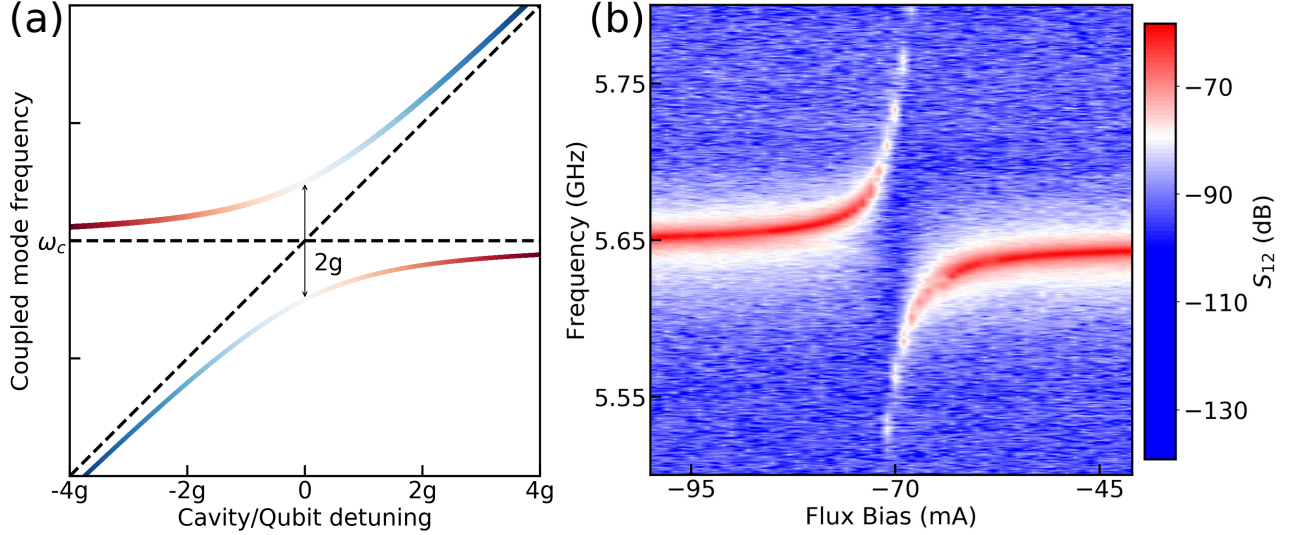


Figure 2.6: **Cavity-qubit avoided crossing:** (a) Energy eigenvalues of the Jaynes-Cummings Hamiltonian plotted as a function of the cavity/qubit detuning Δ for $n = 1$. The bare qubit/cavity energies are plotted as dashed lines. And the JC Hamiltonian eigenstates as colored lines: the more red the line is, the more “cavity-like” the eigenstate is, and the more blue the more “qubit-like”. Near $\Delta = 0$, the eigenstates become equal and orthogonal superpositions of the cavity/qubit, and the eigenvalues undergo an avoided crossing. (b) Spectroscopy of a coupled electromagnetic cavity + tunable transmon qubit system. As the flux bias is tuned, the transmon resonant frequency passes through the (constant) cavity mode, and the two systems undergo an avoided crossing. In this data, $g/(2\pi) \approx 85$ MHz.

2.3.3 The strong dispersive limit and readout in cQED

We now turn our attention to the dispersive regime of the cQED, where the qubit and cavity are detuned by much more than their interaction strength ($g/\Delta \ll 1$) and the eigenstates of the JC Hamiltonian are the “dressed” cavity and qubit states. In this regime, rather than exactly diagonalizing the JC Hamiltonian, it is useful to expand the Hamiltonian in powers of g/Δ . Expanding to second order (see [31] for a derivation), we find

$$\mathbf{H}_{\text{JC,disp}} \approx \hbar \left[\omega_c + \frac{g^2}{\Delta} \sigma_z \right] (\mathbf{a}^\dagger \mathbf{a} + \frac{1}{2}) + \frac{\hbar}{2} \omega_q \sigma_z \quad (2.46)$$

This Hamiltonian looks like the uncoupled cavity/qubit energies, except the cavity energy has been shifted by an amount *that depends on the state of the qubit*. More precisely, the

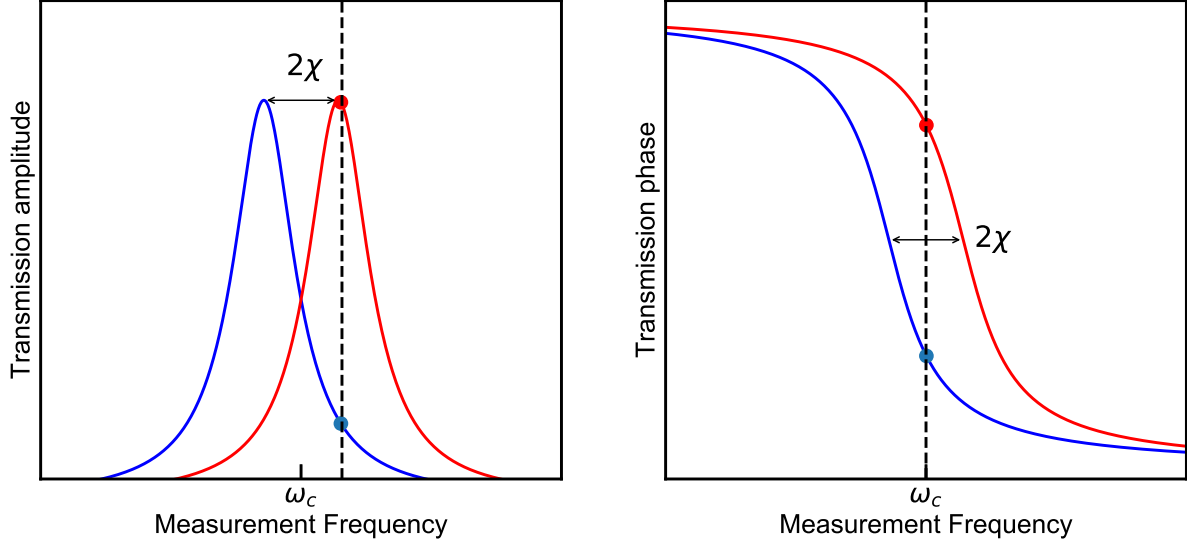


Figure 2.7: **Cavity dispersive shift:** A signal transmitted through an electromagnetic resonator typically takes a Lorentzian shape as a function of frequency, centered on the resonant frequency ω_c with linewidth κ . If the qubit induced dispersive shift $\chi \gtrsim \kappa$, the effective resonant frequency of the cavity $\omega_c \pm \chi$ will shift appreciably and we may distinguish the state of the qubit by measuring the amplitude (a) or phase (b) of the signal transmitted through the cavity. Typically, the phase signal has a higher sensitivity [68, 69].

cavity resonant frequency will be shifted from its bare frequency by $\chi = \pm g^2/\Delta$ if the qubit is in its ground/excited state. If the cavity linewidth κ is small enough such that $\chi \geq \kappa$, the amplitude (or phase) of a near-resonance microwave signal transmitted through (or reflected off) the cavity will change appreciably depending on the state of the qubit, as shown in the Figure 2.7. We also require $\chi \geq \gamma_1$, the decay rate of the qubit, or else the qubit will decay faster than the measurement time $\tau \sim 1/\chi$ required to discern between the two states of the cavity. If these two conditions are met, the signal transmitted through the cavity near the resonant frequency ω_c imparts information about the state of the qubit, and we may use the

transmission through the *cavity* to measure the state of the *qubit*.⁸

With the nonlinearity required to constrain ourselves to a two-state system given by the Josephson junction, and the ability to manipulate and measure this qubit system given by the cQED architecture, incredible progress has been made in using cQED as a basis for quantum information processing. The transmon, first realized in 2007 [34, 38, 72], is now a workhorse device for large number of quantum computing experiments [73] including a proof of principle computation demonstrating the ability of a superconducting quantum processor to perform a computation faster than the world’s most powerful classical supercomputer [54]. These experiments and demonstrations all employ what are fundamentally scaled up versions of the circuit presented in § 2.5: superconducting transmon qubits, controlled and read out via cQED.

⁸In fact, dispersive measurements in cQED are sufficient for executing quantum nondemolition (QND) measurements of the qubit energy eigenbasis. A QND measurement is a “textbook” ideal measurement: the back action of the measurement on the system being measured only projects the system into the observable eigenstate, such that the system could repeatedly be measured producing the same result over and over again [70]. When discussing QND measurements, it is convenient to divide the full device into a “system” with Hamiltonian $\mathbf{H}_{sys}$ which we wish to measure, a “meter” with some Hamiltonian $\mathbf{H}_{meter}$ with which we will measure the system, and an interaction term $\mathbf{H}_{int}$ that describes the interaction between the system and the meter [70, 71]. In this division, $\mathbf{H}_{tot} = \mathbf{H}_{sys} + \mathbf{H}_{meter} + \mathbf{H}_{int}$ and the condition for a QND measurement of some observable $\mathbf{A}$ is that $[\mathbf{H}_{tot}, \mathbf{A}] = 0$. We can make this division easily in Eqn. 2.46, identifying the qubit as the “system”, the bare Harmonic oscillator as the “meter”, and the dispersive shift as the interaction between them. Since $[H_{JC,disp}, \sigma_z] = 0$, a measurement of the cavity will project the qubit into an eigenstate of the Hamiltonian, satisfying the condition for a QND measurement.

Chapter 3

Superconducting qubits in the lab

In our description of superconducting qubits and circuit quantum electrodynamics in Chapter 2, our discussion largely centered around ideal circuits, with difficulties like loss or decoherence only referenced tangentially. Arguably the most important advancement of quantum information science in the past 15 years has been the rapid development of experimental methods and technologies by which we can take these idealisms and bring them into the lab. In this chapter, we tether ourselves back to reality and discuss the devices and experimental setup by which we build and manipulate systems that (at least approximately) are described by the physics of Chapter 2. We will also give a survey of basic experimental methods by which we can characterize these systems.

3.1 The 3D transmon geometry

The main type of qubit employed in this work is the so called “3D transmon” qubit. In the 3D transmon geometry, the qubit and cavity system of cQED are largely the same as in Chapter 2. However rather than a lumped element LC oscillator, the role of the linear electromagnetic oscillator is played by a macroscopically large waveguide cavity resonator, typically a slot milled out of a copper or aluminum box. Fig. 3.1(a) shows a schematic of such a 3D electromagnetic cavity, and Fig. 3.1(b) shows a picture of the bottom half of one of these boxes.

3.1.1 Electromagnetic modes in 3D cavities

If we take the walls of the cavity resonator to be perfectly conducting surfaces (a good approximation for copper, and an excellent approximation for aluminum well below its superconducting transition temperature of 1.2 K), we may approximate the structure as a 3D rectangular box subject to the boundary conditions that the tangential components of the electric field $\vec{E}$ must go to zero at the walls of the box. These boundary conditions force solutions to the electromagnetic wave equation to be built out of *discrete* Fourier components, each of which will have an associated resonance. A box with dimensions L_x, L_y and L_z bounded by conducting walls (see Fig. 3.1) admits classical electromagnetic modes at frequencies [74]

$$\omega_{mnl} = c \sqrt{\left(\frac{m\pi}{L_x}\right)^2 + \left(\frac{n\pi}{L_y}\right)^2 + \left(\frac{l\pi}{L_z}\right)^2} \quad (3.1)$$

where m, n and l are the integer mode numbers, and c is the speed of light in the cavity that depends on the dielectric constant of the material inside of the cavity. If the quality factor of each mode $Q = \omega/\kappa \gg 1$, each mode may be modeled as a noninteracting harmonic oscillator that, when quantized, will act as a single electromagnetic harmonic oscillator that can couple to the qubit. Upon quantizing, we may write down the Hamiltonian of the electromagnetic field in the cavity as a sum of discrete, uncoupled harmonic modes $\mathbf{H}_{EM} = \sum_n \hbar \omega_n \mathbf{a}_n^\dagger \mathbf{a}_n$ [75]. We typically work in the regime $L_x \gtrsim L_z \gg L_y$, such that the lowest frequency, or fundamental, mode is the TE_{101} mode, with $m, l = 1, n = 0$. We choose this mode to be the “cavity” mode from cQED: since $\omega_{01}/(2\pi)$ is typically 4 – 8 GHz, the

cavity mode is generally engineered to be in the 5 – 9 GHz range¹, meaning the long axes of cavity are roughly 2–3 cm. The electric field profile for this mode is plotted in Fig. 3.1(c): to maximize qubit and cavity coupling, the qubit should be as close to the center of the cavity as possible (where the electric field is maximal) and the capacitor pads that provide coupling to the cavity should be aligned with the electric field (to maximize voltage fluctuations across the junction, see § 3.1.2.)

It is reasonable to ask whether or not the assumption that we may ignore the higher modes is valid. The presence of higher modes will renormalize the spectrum and coupling between the qubit and the fundamental mode [78], and may contribute to qubit dephasing (caused by fluctuating mode occupation [76]) and depolarization (via decay into the finite density of states associated with those modes, i.e. Purcell emission [65].) However, there are several reasons these concerns are alleviated in our experiments. If $L_x \approx L_z$, the second lowest frequency modes (TE_{102} and TE_{201}) will be of order $\sqrt{(5/2)}f_0$, see Eqn. 3.1. For a cavity with $f_0 \sim 7$ GHz, this places the next-lowest modes around ~ 11 GHz, rendering coupling to a qubit in the 4-6 GHz range small. Furthermore, the electric field associated with these modes (and the next-next highest T_{202} mode) will be close to zero at the center of the cavity, greatly reducing the coupling to a qubit placed there, see Fig. 3.1 [41]. Additionally, since $\hbar \times 2\pi \times 11 \text{ GHz} > 500 \text{ mK} \gg \sim 10 \text{ mK}$, the dilution refrigerator base temperature, if the measurement lines are properly filtered and the photon bath in these modes is well thermalized, all cavity electromagnetic modes will be in their quantum ground state with all-but unit probability. Note that if the microwave lines that couple to the experiment are

¹We use the fundamental mode of the cavity for several reasons. Since the wavelength of light at microwave frequencies is $\sim$ cm scale, using higher modes would require a rather large resonator for their frequency to be similar the qubit frequency. The fundamental mode of a distributed resonator will also, in general, have the highest quality factor of all the modes [74, 76, 77]. There is also a significant increase in Purcell emission when the qubit frequency ω_{01} is positioned between two cavity resonances, so it is generally advisable to fabricate qubits with a slightly lower frequency than the fundamental cavity mode, see Ref. [65].

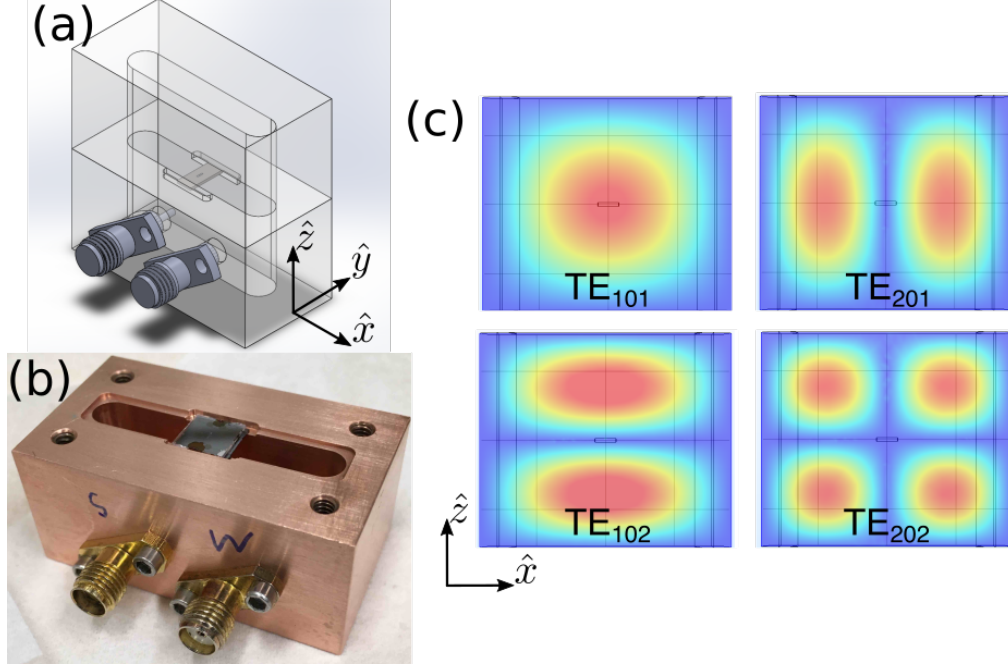


Figure 3.1: **3D electromagnetic cavities:** (a) A CAD model of a typical 3D electromagnetic cavity used as the linear electromagnetic resonator in our lab. The qubit sits on a silicon substrate in the center of the cavity, and we address the cavity via two $50\ \Omega$ lines that antenna couple to the cavity mode. (b) Picture of a copper cavity. The cavity we use is cut in half so we may load and unload qubit devices in the center. (c) Finite element simulation of the magnitude of the electric field for the 4 lowest frequency modes of the cavity. When the qubit is positioned at the center of the cavity, it sits at the antinode of (i.e. is strongly coupled to) the lowest frequency TE_{101} mode, and at a node of (weakly coupled to) the 3 next lowest frequency modes. For the all modes shown here, $\vec{E}$ is parallel to $\hat{y}$.

improperly filtered, this assumption breaks down and the higher frequency modes can drive significant dephasing [76].

In order to manipulate or read out the state of the cavity/qubit system, we must have some sort of external coupling to the cavity resonator. We achieve this coupling with two microwave connectors (typically SMA) terminated to conducting antennas, which extend into the cavity and provide coupling between the microwave lines and the cavity. We quantify the strength of coupling between the cavity and microwave feedlines via the cavity decay rate κ which is the sum of internal (lossy) contributions and external loss into the feedlines

$$\kappa = \kappa_{in} + \kappa_{ext} \quad (3.2)$$

The total decay rate κ is readily extracted from cavity transmission measurements (see section 3.3.1) as the full width at half max (FWHM) of the Lorentzian cavity response [75]. The external coupling κ_{ext} is readily controlled by the length of the input/output antenna. Typically, signal-to-noise is limited by the signal we can collect from the cavity, and not by the power we can apply to the cavity. Therefore we use an asymmetrically coupled cavity, with the input microwave line weakly coupled (short antenna) and the output (measurement) microwave line strongly coupled (long antenna) to maximize cavity decay into the measurement line. Plotted in Fig. 3.2 is a measurement of κ versus the length of the output coupling antenna. In these measurements, the length of the input antenna is fixed and short enough such that, in the absence of the output pin, κ is dominated by internal losses (i.e., $\kappa \approx \kappa_{in}$ for short antenna lengths.)

3.1.2 3D transmon qubits

The large mode volume of 3D microwave cavities presents an intrinsic problem when we consider coupling to a qubit: in order for the cavity-qubit coupling g to be large, the zero point voltage fluctuations² of the cavity V_{ZPF} must efficiently induce voltage fluctuations across the Josephson junction. In a 3D microwave cavity, these voltage fluctuations occur over $\sim$ cm length scales and thus the voltage gradient is small: if the transmon circuit is microscopic/mesoscopic, a large voltage will not build up across the junction, and the

²In Chapter 2 we derived $V_{ZPF} = \omega_c \sqrt{\hbar Z_0 / 2}$ by considering the zero point fluctuations across an LC oscillator. The final result, however, only contains the oscillator frequency and characteristic impedance, and is true for any distributed electromagnetic mode as well. The impedance of a 3D cavity with no dielectric is simply the impedance of free space $Z \approx 377 \Omega$.

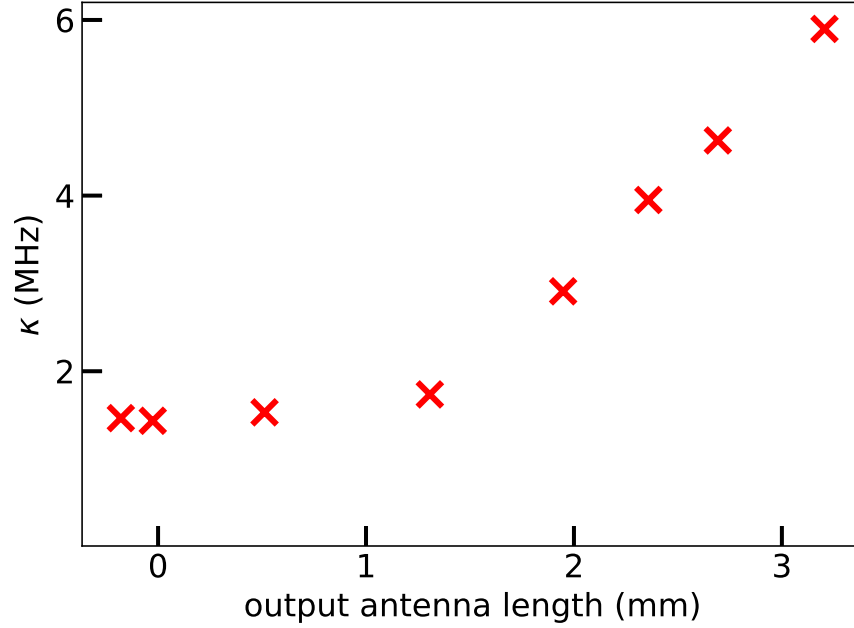


Figure 3.2: **Cavity decay rate κ vs. pin depth:** Measured decay rate κ of a 3D microwave cavity as a function of the output antenna length with the input antenna length fixed. The length is referenced relative to the inner surface of the microwave cavity (and thus may be negative if the pin is recessed from the cavity wall.) These measurements were taken in a polished copper cavity at $T \approx 300$ K.

cavity-qubit coupling will be suppressed.

The solution to this problem is to fabricate a *large* qubit [40]. A 3D transmon (see Fig.3.3) consists of a Josephson junction connected to two coplanar $\sim$ mm scale pads. These pads perform a double purpose: first, they act as the shunt capacitance C_S which brings the qubit into the transmon regime. The “typical” geometry³ shown in Fig.3.3(a), when fabricated on silicon, provides a shunt capacitance $C_S \approx 70$ fF, which gives $E_c \approx 270$ MHz. Second, the pads increase the spatial extent of the qubit, increasing the efficiency with which voltage fluctuations in the cavity induce fluctuations across the junction, and thus increasing

³This geometry is by no means unique: in principle, any geometry with $\sim$ mm spatial extent that provides the desired capacitance will suffice, see Ref. [79] for some examples. In fact, later in this thesis we will discuss some devices with much different geometry than this. The capacitances of these structures may be readily calculated using commercial finite element modeling software, such as COMSOL.

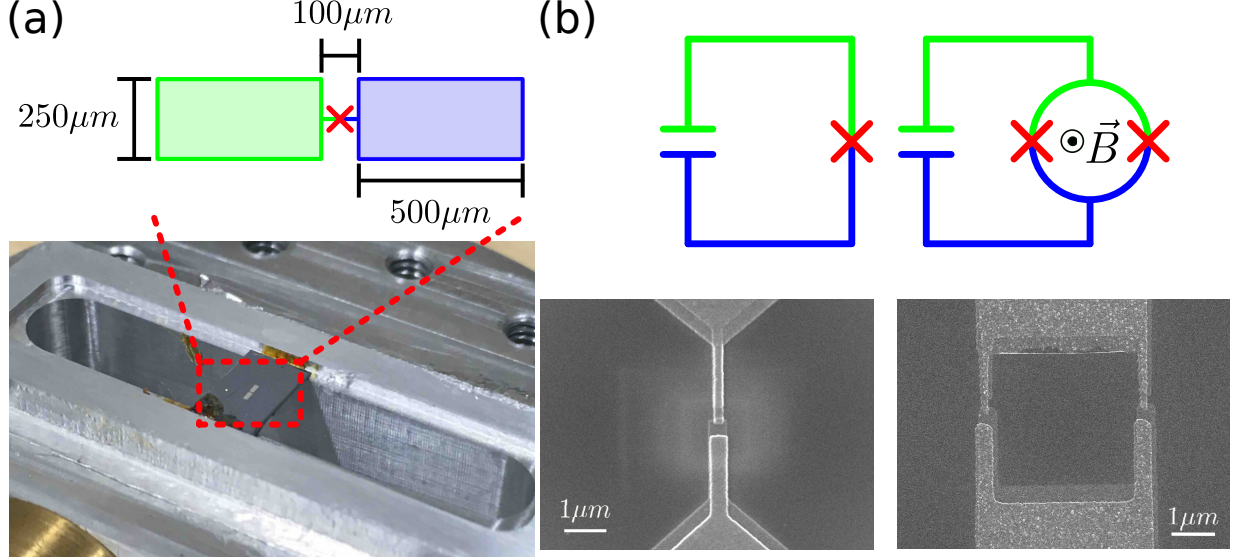


Figure 3.3: **The 3D transmon qubit geometry:** (a) Schematic of a “typical” 3D transmon geometry, inspired by the design in Ref. [40]. The qubit consists of two aluminum pads, typically fabricated on high resistivity silicon, connected by an aluminum/aluminum oxide/aluminum Josephson junction. The two pads simultaneously provide the shunt capacitance (≈ 70 fF for this design) to bring the device into the transmon regime and also allow the qubit to couple strongly to the large volume 3D cavity. Bottom: picture of a 3D transmon as it sits in an aluminum cavity. The pads are easily visible with the naked eye. (b) The two pads are either connected by a single junction (left) or by a “split junction” SQUID loop (right), which allows us to change the effective Josephson energy (and thus the qubit frequency) by applying an external dc magnetic flux. Bottom: scanning electron microscope images of a single junction (left) and SQUID loop (right) fabricated by double-angle evaporation.

the cavity-qubit vacuum Rabi splitting g . This makes for an interesting situation, where a quantum system that hosts a collective, coherent degree of freedom we can manipulate is visible to the naked eye!

While the whole qubit is $\sim$ mm scale, the Josephson junction that provides the nonlinearity is only $\sim 100 \times 100$ nm². Junctions are fabricated using the Dolan bridge, double-angle evaporation process [80]; a full description of the fabrication process is given in Appendix C. When designing an experiment, perhaps the most critical parameter to consider is the transmon $|0\rangle \rightarrow |1\rangle$ transition frequency $\hbar\omega_{01} = \sqrt{8E_J E_C} - E_C$. As mentioned above,

$E_C = e^2/2C$ is controlled by the geometry of the (macroscopic) pads, and is thus fairly easy to design and fabricate with a high degree of accuracy. The Josephson energy, however, is more subtle: it is proportional to the critical current of the junction $E_J = \hbar I_c/2e$, and depends on the thickness of the oxide layer and the area of the junction. Since we have no way of readily measuring the oxide layer thickness, and the junction is small enough that typical variations in the fabrication may cause appreciable changes in the junction area, it is difficult to *a priori* predict E_J for a given design. It is, however, possible to infer the critical current of the junction before cooling the device down: the Ambegaokar-Baratoff relationship predicts that the critical current of a Josephson junction is inversely proportional to the normal-state resistance of the junction R_n [21]

$$I_c = \frac{\pi \Delta_{\text{BCS}}(0)}{2eR_n} \quad (3.3)$$

where $\Delta_{\text{BCS}}(0) = 170 \text{ } \mu\text{eV}$ is the $T = 0$ superconducting gap of aluminum. We may therefore determine the critical current, and thus E_J , of the junction simply by measuring the resistance of the junction at room temperature⁴. These tests are crucial for quick feedback in the fabrication process, and also for verifying that the junctions are working, since small junctions are notoriously easy to destroy with static discharge.

Even with normal state resistance testing at our disposal, variations in the junction fabrication process make it difficult to obtain a desired E_J with high precision. We normally have variations in E_J of $\sim 10\%$ between nominally identical junctions fabricated on the same silicon wafer. State-of-the-art industrial processes have yielded junction critical current variations approaching 0.5-1% [81, 82], however these processes are untenable with

⁴We measure the normal state resistance of these junctions using a home built “probe station”. Instructions on how to use the probe station may be found in the lab shared drive. A shortcut I find useful when probing qubits is $E_J/h \approx 134 \text{ GHz} / (R_n \text{ in k}\Omega)$

our fabrication capabilities. There is, however, a simple way to tune the critical current *in situ*: we can modify the design of the qubit slightly, replacing the single junction with two junctions in parallel, enclosing some area A and forming a loop (see Fig. 3.3(b).) For the purpose of this dissertation this geometry, called a superconducting quantum interference device (SQUID), acts like a *single* junction with a critical current that may be modulated by threading the loop with a magnetic flux $\Phi = \vec{B} \cdot \vec{A}$. For a symmetric SQUID, consisting of two junctions each with critical current I_c , the total critical current of the SQUID loop is given by [21, 83]

$$I_S = 2I_c \left| \cos \left(\frac{\Phi}{\Phi_0} \right) \right| \quad (3.4)$$

where, once again, $\Phi_0 = \hbar/2e$ is the reduced superconducting flux quantum. The effective Josephson energy E_J is modified proportionally, allowing us to change E_J (and thus the qubit frequency) by applying a dc magnetic field. We apply this magnetic field by driving a current through a coil of superconducting wire that winds around the cavity, as shown in Fig. 3.4. Note that, since superconductors expel magnetic fields, in order to “flux tune” the qubit the cavity must be made of a metal that is not superconducting at 10 mK: copper is by far the best option, since it is an excellent thermal conductor at low temperatures, and low-magnetic impurity oxygen free high conductivity (OFHC) copper is readily available at reasonable prices.

While the extra handle of *in situ* E_J tunability is experimentally handy, it also opens up another noise channel, by which the transmon may decohere, as flux noise may readily modify the qubit frequency. To mitigate flux noise, the area A of the SQUID loop should be kept as small as possible, while still being large enough to tune the qubit by passing a

reasonable current through the coil. In our lab, we find that $A = 16 \mu\text{m}^2$ is sufficient such that a 30 turn coil can provide a flux quantum to the loop by passing ≈ 130 mA through the coil.

3.2 Filtering and shielding

As stated in Chapter 2, the nonlinearity of the Josephson Junction is a necessary but insufficient requirement for operating the cavity-qubit system in the quantum regime: we also need both the cavity and the qubit to be initialized with high probability in the ground state, and for few thermal excitations to decohere the system. This requirement imposes the constraint $\hbar\omega_{01}, \hbar\omega_c \gg k_bT$: for a sense of scale, $5 \text{ GHz} \approx 240 \text{ mK}$. This means that the temperature at which these experiments operate at must be of order 10's of mK. Nominally, these temperatures are straightforward to achieve in commercial dilution refrigerators: all superconducting qubit experiments in our lab are conducted on a BlueFors LD-400 cryogen-free dilution refrigerator with a nominal base temperature of 7 mK. However, the temperature read out by the thermometer on the mixing chamber plate of a dilution fridge does not necessarily reflect the temperature of a thermal bath coupled to the cavity-qubit system. Spurious excitations in the cavity/qubit system, as well as external noise that couples to a qubit, are a major source of decoherence.

Spurious excitations and decohering noise may come from the ambient environment of the experiment, or through the microwave feedlines with which we manipulate/measure the system. Thus, quite a bit of work has gone into optimizing the shielding of the cavity/qubit system, as well as filtering on the input/output lines. The setup in our lab, detailed below, is fairly standard, though there is some room for improvement.

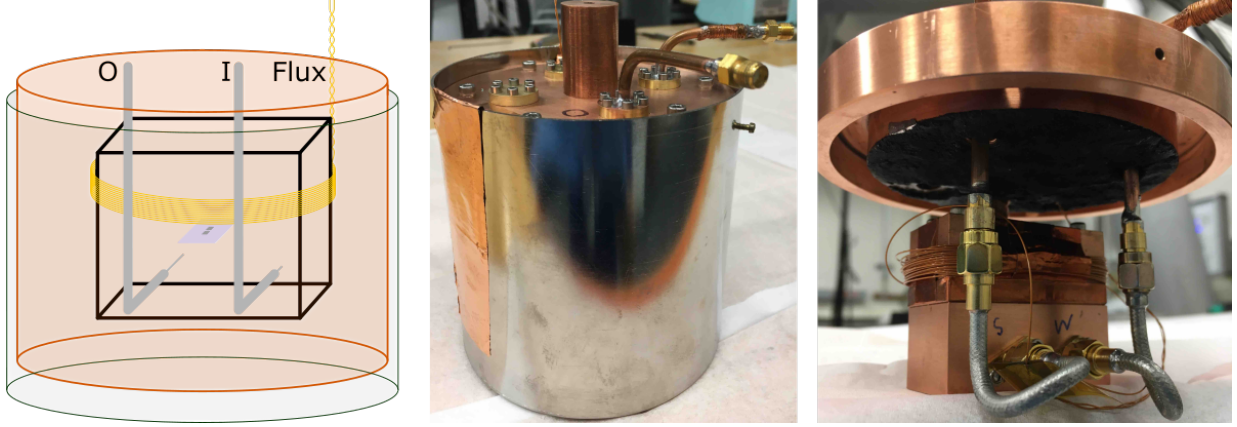


Figure 3.4: **Cryogenic shielding for superconducting qubit experiments:** Left: schematic of the 3D cavity as it sits in the copper shielding box and Cryoperm magnetic shielding. Center: Fully assembled, shielded cavity on the lab bench. Visible is the Cryoperm shielding (silver), the hermetic copper box and the input/output coaxial cables. Right: partially disassembled box, displaying how the 3D cavity is connected to the microwave/flux lines. Also visible in this picture is the IR absorbing material (black surface on the bottom of the lid) that coats the inside of the copper box.

3.2.1 Cryogenic shielding

It has been shown that superconducting qubit devices may be susceptible to stray infrared radiation and microwave inside the cryostat [84, 85]. To shield the experiment from stray photons in the cryostat, we place the cavity inside a copper “box”, which is hermetically sealed with indium o-rings to prevent thermal radiation from leaking in through the seams (see Fig. 3.4.) In newer iterations of the experimental setup, the input/output lines, as well as the flux tuning line, are fed through the lid of the box via brass flanges, which are also hermetically sealed with indium o-rings. The inside walls of the box are coated with a home-made infrared absorber dubbed “Berkeley black” [86] consisting of Stycast 2850FT epoxy (73% by weight⁵), carbon lampblack (7% by weight) and 175 nm diameter glass beads (20% by weight.) This IR absorbing material is visible on the lid of the box in the right-most

⁵The reference cited calls for Catylist 24LV, however we used Catylist 9, and changed the ratios such that the total weight of the epoxy added up to 73%.

picture in Fig. 3.4.

In addition to shielding from stray IR, superconducting qubit devices (especially split-junction tunable devices) are susceptible to magnetic flux noise. To avoid noise from stray magnetic fields in the lab, and to cool the device in as low a magnetic field as possible to avoid trapped vortices, we further shield the cavity/box with a high magnetic permeability metal (trade name “Cryoperm”) as seen in the middle panel of Fig. 3.4.

3.2.2 Input/output line filtering

An inescapable consequence of the input/output pins needed to control/read out the state of the system is added noise and loss. In fact, a completely lossless system would be of little interest to us: we *want* photons in the microwave cavity to decay into the transmission line connected to the output pin, from which we may infer the state of the system! While we must incur this penalty to be able to manipulate the system, we would like to filter away as much unwanted noise as possible. As such, we extensively filter the input and output microwave lines; this filtering is detailed in Fig. 3.5.

The input/output microwave lines may be a major source of noise in our experiments. Without proper filtering, the coax cables may act as a waveguide that brings high-frequency electromagnetic radiation from higher temperatures to mixing chamber plate at ~ 10 mK. Radiation in the 5-20 GHz range may populate the lower frequency modes of the cavity, driving dephasing [76, 87, 88], while higher frequency (> 100 GHz) radiation may break Cooper pairs in the qubit and cause depolarization [42]. Furthermore, the coaxial structure of the microwave lines presents a problem: it is difficult to thermalize the inner conductor of the coax, since it is galvanically disconnected from the outer (ground) conductor, and thus

cannot transfer heat to the dilution refrigerator through conduction electrons⁶ [87, 90].

To alleviate these concerns, at every stage of the cryostat there is an “attenuator” on both the input and output lines. These are either real attenuators, in the form of voltage dividers that both attenuate signals from higher stages of the cryostat and provides galvanic (and thus thermal) contact to ground, or so called “0 dB attenuators”, which are micro-strip heat exchangers in attenuator packaging that provide thermal connection to ground without significant attenuation. Note that we use real attenuators on the input line, and thus by necessity must attenuate the control signals we send into the fridge. Attenuating the signal also causes resistive heating: we therefore distribute the attenuators across the different cryostat stages such that the majority of the resistive dissipation occurs at higher temperature stages (with more cooling power) while the qubit is shielded from potentially harmful thermal radiation from the 4K and still (~ 800 mK) plates [91].

In addition to attenuation, we also use two types of lowpass filters on both the input and output lines to further attenuate high frequency radiation. We use a commercial lumped-element low-pass filter to attenuate radiation above the operating frequency of the cavity/qubit system. We have observed that these filters greatly increase T_2 of the qubit, likely because they filter away radiation that populates the higher-frequency modes of the cavity, which is known to be a major source of dephasing [76]. We also use home-made filters (based on the design in Ref. [92]) that utilize a lossy-magnetic epoxy (trademarked under the name “Eccosorb”) to attenuate high-frequency signals. Since the Eccosorb filters do not contain lumped element components, they do not have detectible re-entrant modes at higher frequencies, and have been shown to strongly attenuate signals up to 100 GHz (i.e. photons

⁶Thermal conduction from phonons in solid insulators becomes vanishingly small at dilution refrigerator temperatures [89]

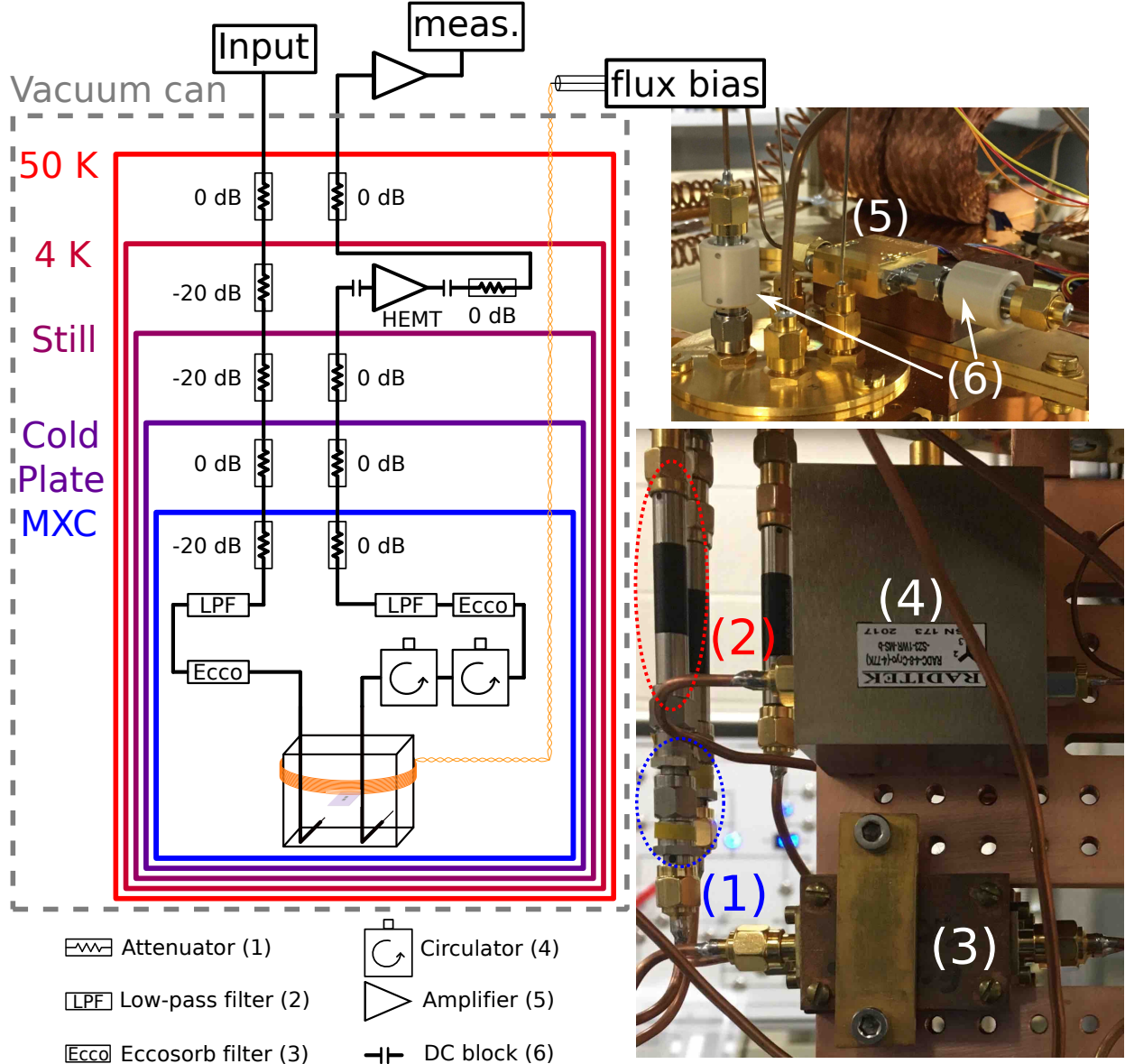


Figure 3.5: **Filtering and amplification of the input-output lines:** Left: schematic diagram detailing the filtering and amplification along the input/output lines of the superconducting qubit setup. In this schematic, each component is thermalized to the cryostat plate corresponding to the box the component resides in (for example, the HEMT amplifier is thermalized to the 4K stage.) Right: pictures of one of each component as they reside on the dilution refrigerator. The top picture, focusing in on the HEMT amplifier, is taken at the 4K stage, while the bottom picture is taken below the mixing chamber plate (MXC.) Note that the amplifier at room temperature is a different model than the one shown (see Table 3.1.)

with enough energy to break Cooper pairs in aluminum.)

While we may attenuate the input lines to protect the cavity/qubit system, the signals we would like to measure are rather weak, and attenuating them by ~ 60 dB would render them undetectable. To protect the system from radiation coming down the output line without attenuating the desired signal, we use a series of microwave circulators. A circulator is a three port device which uses a ferrite core to break reciprocity: an ideal circulator would allow signals entering one port to propagate only to the port counter-clockwise (i.e., if ports were labeled 1, 2 and 3 in a counterclockwise fashion, signals from port 1 could propagate to port 2 but not port 3.) In reality, the circulators we use provide ≈ 30 dB of isolation from backwards-propagating signal, so we use two circulators in series to achieve robust isolation. The unused port of the circulator is capped off with a $50\ \Omega$ terminator.

3.2.3 Amplification

In addition to protecting the cavity/qubit system from decohering noise, we also need to amplify the signal that encodes the state of the system. Amplifying a signal *also* inevitably adds noise to the signal, on top of amplifying any noise that was previously being carried along with the unamplified signal [93]. To understand how the noise added by an amplifier is quantified, it is useful to recall the Johnson-Nyquist relation for thermal noise: the power spectral density (PSD) of noise emitted into a transmission line by an impedance matched load is proportional to the temperature of the load [93, 94, 71]

$$S[\omega] = k_b T \tag{3.5}$$

The noise PSD is defined as the power per unit bandwidth of the noise: if we had a

detector with bandwidth B that detected signals only between ω_a and $\omega_a + B$, the total power of the noise that would be detected would be $P = \int_{\omega_a}^{\omega_a+B} d\omega S[\omega] = k_b T B$. Importantly, Johnson-Nyquist noise is “white noise”, i.e. it has no frequency dependence⁷. This conveniently allows us to define the noise temperature of an amplifier: imagine an amplifier with power gain G impedance matched to a transmission line that carries a signal with noise at temperature T_s . If the amplifier added no additional noise, the noise PSD of the output signal would simply be given by

$$S_{\text{out}}[\omega] = G k_b T_s \quad (3.6)$$

However, it can be shown [93] an amplifier must add additional noise to the signal: a real amplifier with *noise temperature* T_n instead outputs a signal with a noise PSD given by

$$S_{\text{out}}[\omega] = G k_b (T_s + T_n) \quad (3.7)$$

Thus, amplifiers are typically characterized by their power gain G and their noise temperature T_n . What happens, however, if we must amplify the signal again? We will, of course, pass the signal through another amplifier, which will *also* add noise to the system: the twice-amplified noise PSD will be given by

$$S_{\text{out},2}[\omega] = G_2 k_b (G_1 (T_s + T_{n1}) + T_{n2}) = G_{\text{chain}} k_b (T_s + T_{\text{chain}}) \quad (3.8)$$

⁷This model breaks down at low temperature and high frequencies, where $\hbar\omega$ becomes comparable to $k_b T$. A full quantum treatment of the noise of the load yields the result $S[\omega] = \hbar\omega/2 \coth(\hbar\omega/2k_b T)$ [93]. At high temperatures, this approaches the Johnson-Nyquist result, however as $T \rightarrow 0$ it bottoms out at $\hbar\omega/2$, the zero point fluctuations of the electromagnetic field at ω . Using frequency and temperature scales that are common in this thesis yields a situation where the Johnson-Nyquist theory breaks down: A 5 GHz signal coming from the MXC plate of the cryostat, held at 10 mK, will have an “effective noise temperature” $T_{\text{eff}} = S[\omega]/k_b = \hbar\omega/2k_b \approx 125$ mK, much hotter than the physical temperature of the load.

where $G_{\text{chain}} = G_1 G_2$ is the total amplification of the chain of amplifiers, and the effective noise temperature of the amplifier chain is given by

$$T_{\text{n,chain}} = T_{n1} + \frac{T_{n2}}{G_1} \quad (3.9)$$

We could in principle⁸ amplify the signal ad nauseam by stringing together more amplifiers: it is relatively easy to see that, for a chain of m amplifiers, the gain of the chain will be $G_1 G_2 \dots G_n$ and the effective noise temperature of the chain will be

$$T_{\text{n,chain}} = T_{n1} + \frac{T_{n2}}{G_1} + \frac{T_{n3}}{G_1 G_2} + \dots \quad (3.10)$$

This relationship tells us that for a chain of amplifiers that have large power gain G_i , the effective noise temperature of the chain (and, thus, overall noise added by the amplifiers) is usually dominated by the *first* amplifier in the chain. Thus, the first stage of amplification is by far the most important when it comes to optimizing signal to noise. The first stage of amplification on our experimental setup is provided by a high electron mobility transistor (HEMT) amplifier at the 4K stage, with a gain of ≈ 36 dB (which corresponds to an absolute power gain $G \sim 4000$) and a noise temperature of ≈ 3.5 K at typical signal frequencies. The second amplifier (at room temperature) has a noise temperature of ~ 110 K, indicating that we are, in fact, in the limit where the amplification noise is dominated by the first amplifier.

3.2.4 Possible improvements

As stated above, the experimental setup in our lab is fairly standard for supeconducting qubit experiments, and is sufficient for the purposes of this thesis. There is, however, ample

⁸In practice we really, really couldn't.

Component description	Supplier/manufacturer	Part number
Cryogenic attenuator (- XX dBm)	XMA Corporation	2082-6418- XX -CRYO
Cryogenic compatible inner/outer dc block	Richardson RFPD/API Tech	8039
50 Ω terminator	XMA Corporation	2001-6113-00
Lumped element low-pass filter	K&L Microwave	6L250-00088
Circulator	Raditek	RADC-4-8-Cryo-4-77K- S23-1WR-MS-b
HEMT amplifier	Low Noise Factory	LNF-LNC03_14A
IQ mixer (operating frequency XX - YY GHz)	Marki Microwave	IQ- XXYY LXP
Room temperature RF-amplifier	Minicircuits	ZX60-83LN-S+
Splitter/Combiner	Minicircuits	ZFSC-2-10G
dc amplifier board	Texas Instruments	DEM-OPA84XU
dc op-amp	Texas Instruments	OPA847
Quasi-dc low-pass filter	Minicircuits	BLP-21.4+

Table 3.1: RF circuitry components for the circuit quantum electrodynamics experimental setup, along with part numbers and suppliers. For some components, we use multiple variants of the same component. For these cases, the variable parameter is marked in **bold**, and the corresponding change in the part number is matched to the variable in the component description. Some of the specific products here are no longer manufactured, however as of writing this, every manufacturer listed either still makes the part listed or makes an updated version of it.

room for improvement on the filtering/shielding in future experiments. Here, we outline some known weaknesses in the current setup, as well as some recent literature that suggests improvements that could be made.

One place for immediate improvement is in the magnetic shielding: when doing experiments using flux tunable transmons, we can see shifts in the qubit frequency when the magnet on the other dilution refrigerator in the lab (roughly 3-4 m away) is varied. One way the magnetic shielding could be improved is by using an superconducting shield (such as aluminum) in addition to the Cryoperm shield [95]. There is also a push in some groups to remove *all* magnetic material from the experimental stage: this would require replacing all stainless steel screws with brass screws. Additionally, recent experiments [42] have also indicated that the concentration of superconducting quasiparticles may be significantly reduced if an additional set of Eccosorb filters on the input/output lines are placed *inside* of the hermetic copper box.

On the other hand, recent experiments have also shown that quasiparticle bursts in superconducting islands could be caused by long time-scale thermalization in the cryostat [96]. The proposed improvements outlined above largely involve adding more superconducting material (which acts as a thermal break, as the superconducting condensate doesn't carry entropy) or insulators, which, without electronic heat conduction, cool inefficiently at dilution refrigerator temperatures. If long timescale thermalization is found to be a limiting source of decoherence once other sources are efficiently eliminated, it may be of great interest to optimize the thermalization of many individual parts of the experimental setup, or to reduce the thermal load of the experiment by judiciously choosing design geometries/materials with a low specific heat when possible.

3.3 Spectroscopic measurements

We now turn to a survey of some basic experimental techniques used in our lab. These experimental techniques largely fall into two categories: spectroscopic experiments, where we experimentally identify energy eigenstates of the system, and time-resolved measurements, where we apply pulses to control the system Hamiltonian as a function of time. We begin with a description of several common spectroscopic measurement techniques, from which we extract the relevant Hamiltonian parameters such as the cavity frequency ω_c or the qubit $|0\rangle \rightarrow |1\rangle$ transition frequency ω_{01} .

3.3.1 “Punchout”

The very first experiment one typically does upon cooling down a new sample is verifying that a functioning qubit is, in fact, coupled to the microwave cavity. We typically accomplish this by measuring the transmission through the cavity at different driving powers. At very low power, the cavity transmission amplitude is well described by a Lorentzian centered on the dispersively shifted cavity frequency $\omega_c \pm \chi$ (depending on whether the ω_{01} is lower or higher than ω_c .) As the drive power increases, a semiclassical model of the system predicts a region of bistability with two solutions that destructively interfere, creating a “dark state” response, followed by the abrupt onset of a Lorentzian response at the bare cavity frequency ω_c at some critical power where the bistable region vanishes [97, 98, 99].

The data plotted in Fig (3.6), taken in our lab by monitoring the transmission through the microwave circuit using an Agilent N5230A vector network analyzer (VNA), agree with the transmission predicted by this model very well. We colloquially refer to the onset of the bright state (occurring at $f \approx 5.775$ GHz and $P_{\text{VNA}} \approx -8$ dBm in Fig. 3.6(b)) as “punchout”:

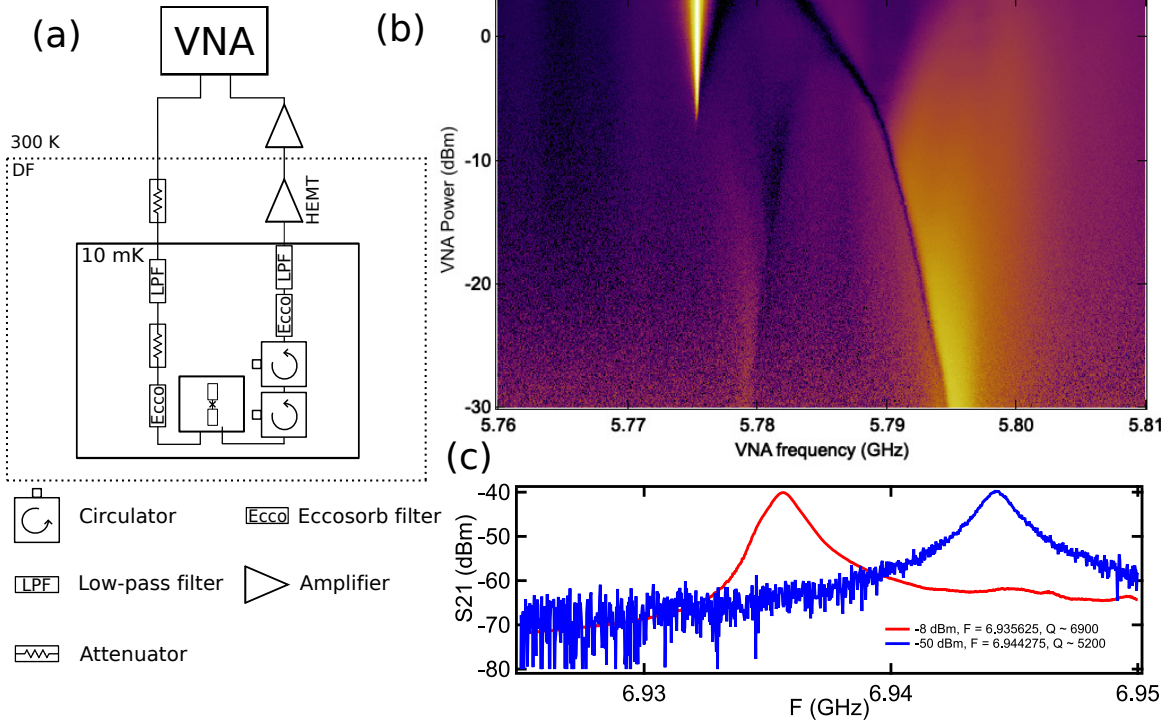


Figure 3.6: **Qubit “punchout”**: (a) Schematic of the measurement setup employed when doing a “punchout” measurement. (b) Transmission through a 3D cavity plotted a function of both VNA power and cavity frequency. At low power, we observe a Lorentzian spectrum centered at the dressed cavity frequency $\omega_c + \chi$. At intermediate powers, the transmission at the dressed frequency is suppressed and a dark state forms in agreement with theory [97, 98]. At high powers, we observe an abrupt onset of a high transmission Lorentzian peak at the bare cavity frequency ω_c . (c) A more typical “punchout” measurement, where we only record the transmission at low and high powers to extract ω_c , χ and the quality factors at these powers. Note that this data corresponds to a different cavity than the data presented in (b).

the high power signal “punches out” the qubit influence on the cavity transmission, and the cavity responds at its classical geometric resonant frequency with $\chi = 0$. A back-of-the-envelope calculation indicates that the power required to punch out the qubit corresponds to a cavity photon occupation of $\sim 10^5$, which agrees well with the number quoted in Ref. [99].

In a typical experiment, rather than mapping out the full cavity response as a function of power, one only measures the low and high power transmissions (see Fig. 3.6(c)). This yields

the cavity resonant frequency ω_c , the dispersive shift χ and the low and high power quality factors of the cavity, which may be extracted from a Lorentzian fit to the data or, more quickly, $Q^{-1} = 2(\omega_0 - \omega_{P/2})/\omega_0$ where $\omega_{P/2}$ is the frequency at which half the maximum power is transmitted.

3.3.2 Two-tone spectroscopy

Upon verifying that a functional qubit is in the cavity, we need to find its $|0\rangle$ to $|1\rangle$ transition frequency ω_{01} . We generally accomplish this through two-tone spectroscopy [100], which takes advantage of the qubit state dependent dispersive shift imparted on the cavity frequency. We use the VNA to continuously monitor the transmission through the cavity at the dressed cavity frequency $\omega_c \pm \chi$. At the same time, we sweep the frequency of another tone impinging on the cavity. When this variable frequency tone matches ω_{01} , we Rabi drive the qubit, changing its average excited state population from $P_1 \approx 0$ to $P_1 \approx 0.5$. Since two-tone measurements typically take place on $\sim$ second time scales while Rabi drive frequencies are typically in the $\sim$ MHz range, the Rabi drive manifests itself as a splitting of the cavity transmission into two peaks (one dispersively shifted by the qubit in the $|1\rangle$ state, and one dispersively shifted by the qubit in the $|0\rangle$ state.) This causes the previously on-resonance transmission measurement to drop in magnitude. Thus, we measure ω_{01} by recording the frequency at which we see a dip in the cavity transmission.

There are several things one should keep in mind while doing a two-tone measurement:

- The linewidth of the qubit is ultimately limited by T_2^{-1} , however width of the transmission dip may be power broadened by either the cavity or qubit drive tones [68, 101].

To obtain as accurate data as possible, one should use a small power for both tones

3.4 Time-resolved measurements

Spectroscopic measurements are generally good diagnostic tools, however time-resolved measurements are a far more powerful technique for manipulating and characterizing the system. In a time-resolved measurement, we drive the system for a fixed amount of time, let it evolve freely for a fixed amount of time, and then at the end measure the state of the system to see how it has evolved. In this manner, we can investigate how different components of the system interact with each other, or how each subsystem interacts with its environment. This section will describe the general measurement scheme used for time-resolved measurements in this thesis, and then provide a (by no means exhaustive) sample of common time-resolved measurement protocols.

3.4.1 General time-resolved setup

The heart of the time-resolved measurement setup is the Tabor WX2184C arbitrary waveform generator (hereon referred to as “the Tabor”.) The Tabor has four outputs, each of which may output a sequence of arbitrary waveforms with a resolution of 1 GS/s. We mix these outputs with continuous wave microwave tones to create the GHz pulses that manipulate/read out the system. To initiate the start of a waveform, we trigger the Tabor with a voltage pulse from the SRS DG535 digital delay generator. The Tabor outputs the waveform, steps to the next waveform in the sequence, and then idles until the next trigger pulse. By this method we may step through sequences of waveforms, which is necessary if we want to make measurements as a function of the pulse parameters (delay time, amplitude, etc...) The Tabor also has lower voltage resolution “marker” outputs which may be used to trigger other components in the circuit at a specified time during the pulse sequence. We use the

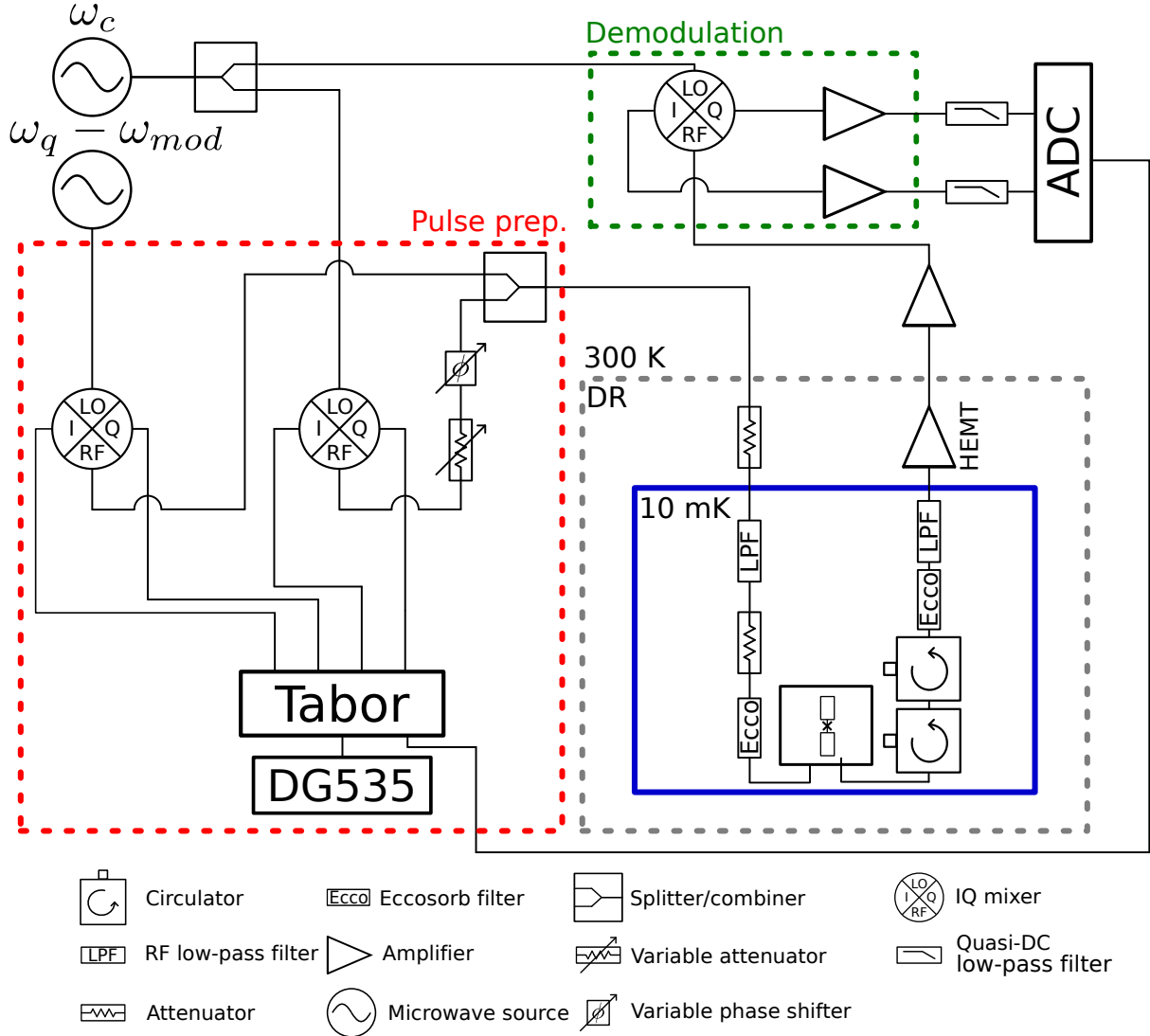


Figure 3.8: **General time resolve measurement setup:** A schematic diagram of the most general measurement setup employed when doing time-resolved measurements in this thesis.

marker output to trigger the analog to digital converter (ADC) to start recording data.

3.4.2 Coherent control

As alluded to in Chapter 2, the conduit by which we manipulate delicate quantum systems is coherent electromagnetic radiation: a driving electric field oscillating at $\omega_d \sim \omega_{01}$ adds a

new term into the qubit Hamiltonian that we have control over. To see this, let's rewrite the Hamiltonian of a transmon qubit in the presence of a (classical) electric field that we derived⁹ in § 2.3.1:

$$\mathbf{H} = \left[\hbar\omega_q(\mathbf{a}_q^\dagger\mathbf{a}_q + \frac{1}{2}) + \mathcal{O}(\phi^4) \right] + \epsilon(t)(\mathbf{a}_q + \mathbf{a}_q^\dagger) \quad (3.11)$$

where, as a reminder, $\mathbf{a}_q^\dagger, \mathbf{a}_q$ are the qubit energy level raising and lowering operators, and $\mathcal{O}(\phi^4)$ contains the nonlinearity of the Josephson junction that allows us to operate the transmon circuit as a qubit. Recall also that $\epsilon(t)$ is a term that encapsulates both the (classical) applied voltage $V(t)$ and the coupling between the transmon and the voltage source. To model a coherent field impinging on the qubit, we will take the applied voltage to be a constant-amplitude sinusoidal drive at frequency ω_d , i.e. $V(t) = V_0 \cos(\omega_d t + \phi)$. We can simplify our lives by encapsulating the amplitude of the sinusoid V_0 along with the coupling in a single constant $A \propto V_0$, and writing the Hamiltonian as:

$$\mathbf{H} = \left[\hbar\omega_q(\mathbf{a}_q^\dagger\mathbf{a}_q + \frac{1}{2}) + \mathcal{O}(\phi^4) \right] + A \cos(\omega_d t + \phi)(\mathbf{a}_q + \mathbf{a}_q^\dagger) \quad (3.12)$$

We now once again take the qubit approximation, saying the nonlinearity is large enough that we may constrain ourselves to the $\{|0\rangle, |1\rangle\}$ manifold¹⁰, and switch from the language of raising and lowering operators to the language of spins by sending $\mathbf{a}_p^\dagger$ to $\boldsymbol{\sigma}_+$ and $\mathbf{a}_p$ to $\boldsymbol{\sigma}_-$. In this approximation, the above Hamiltonian takes the familiar form of the “Rabi problem”

⁹In a footnote in § 2.3.1, I mentioned that the choice of the drive term $\propto (\mathbf{a}_q - \mathbf{a}_q^\dagger)$ vs $\propto (\mathbf{a}_q + \mathbf{a}_q^\dagger)$ is simply a matter of convention. To be consistent with the literature, in this section I have opted to work with the dependence in the form $\propto (\mathbf{a}_q + \mathbf{a}_q^\dagger)$: in the two-level approximation this is equivalent to changing the driving term in the Hamiltonian from $\boldsymbol{\sigma}_y$ to $\boldsymbol{\sigma}_x$, which, as can be seen from Eqn. 3.17, we may accomplish by adding in an arbitrary phase shift to the drive.

¹⁰While we use $|0\rangle$ and $|1\rangle$ as our canonical example, it is important to remember that the results presented here are valid for any two states with sufficient nonlinearity to prevent leakage into adjacent states. For example, we may just as well Rabi drive between the $|1\rangle$ and $|2\rangle$ states.

$$\mathbf{H}_{\text{Rabi}} = -\frac{\hbar\omega_{01}}{2}\boldsymbol{\sigma}_z - A \cos(\omega_d t + \phi)\boldsymbol{\sigma}_x \quad (3.13)$$

If $|\omega_d - \omega_{01}| \ll \omega_{01}$, we may again invoke the rotating wave approximation and solving for the time dependence of the state probability amplitudes using first-order time-dependent perturbation theory. Doing so yields the famous Rabi oscillation result [36], where $|0\rangle$ and $|1\rangle$ state amplitudes C_0 and C_1 oscillate back and forth at a rate $\Omega_R = \sqrt{(\omega_d - \omega_{01})^2 + A^2/\hbar^2}$ that depends both on the drive amplitude and the detuning between the drive frequency ω_d and the qubit frequency ω_{01} (see Fig 3.9 (a-b).) More succinctly, the excited state probability as a function of time is

$$P_1(t) = |C_1(t)|^2 = \frac{A^2}{(\hbar\Omega_R)^2} \sin^2(\Omega_R t/2), \quad (3.14)$$

It is more instructive to work with the Hamiltonian in the frame *rotating with the drive*. We change basis by applying a unitary operator that “undoes” the rotation caused by the drive

$$\mathbf{U} = e^{-i\frac{\omega_d}{2}\boldsymbol{\sigma}_z t} \quad (3.15)$$

To transform the Hamiltonian, we apply this unitary to the time-dependent Schrödinger equation and solve for an effective Hamiltonian in the rotating frame

$$\begin{aligned} \mathbf{U} i\hbar \frac{\partial}{\partial t} |\psi\rangle &= \mathbf{U} \mathbf{H} |\psi\rangle \\ i\hbar \left(\frac{\partial}{\partial t} (\mathbf{U} |\psi\rangle) - \frac{\partial \mathbf{U}}{\partial t} |\psi\rangle \right) &= \mathbf{U} \mathbf{H} \mathbf{U}^\dagger \mathbf{U} |\psi\rangle \\ i\hbar \frac{\partial}{\partial t} \mathbf{U} |\psi\rangle &= \left[\mathbf{U} \mathbf{H} \mathbf{U}^\dagger + i\hbar \frac{\partial \mathbf{U}}{\partial t} \mathbf{U}^\dagger \right] \mathbf{U} |\psi\rangle \\ \tilde{\mathbf{H}} &= \mathbf{U} \mathbf{H} \mathbf{U}^\dagger + i\hbar \frac{\partial \mathbf{U}}{\partial t} \mathbf{U}^\dagger \end{aligned} \quad (3.16)$$

If we plug our Hamiltonian 3.13 into the above equation, after doing some algebra and invoking the rotating wave approximation, we find

$$\tilde{\mathbf{H}}_{\text{Rabi}} = -\hbar \frac{(\omega_{01} - \omega_d)}{2} \boldsymbol{\sigma}_z - \frac{A}{2} (\cos(\phi) \boldsymbol{\sigma}_x + \sin(\phi) \boldsymbol{\sigma}_y) \quad (3.17)$$

The implication of Eqn. 3.17 is that, in the frame rotating with the drive, we may control the quantization axis by the qubit-drive detuning $\omega_{01} - \omega_d$ and the phase of the drive ϕ (several different example conditions are shown in in Fig. 3.9(c-e).) Thus, it is important for us to have both frequency and phase sensitive control of the drive tone.

Returning to the measurement setup shown in Fig. 3.8, we employ microwave frequency IQ mixers to mix the low-frequency modulation signal from the Tabor with the high frequency continuous-wave tone from the microwave source. An IQ mixer is a 4-port of mixer that, given a tone into the Local Oscillator (LO) port $A_1 \cos(\omega t)$ and signals $I(t)$ and $Q(t)$ into the I and Q ports respectively, outputs (ideally)

$$RF = A_1 (I(t) \cos(\omega t) - Q(t) \sin(\omega t)) \quad (3.18)$$

which is visualized in Fig. 3.10. Here, I stands for “in phase” and Q stands for “quadrature.” Modulating the signal using an IQ mixer is particularly important, since it allows us to control the phase of the microwave qubit control pulse. If we choose $I(t) = A_2 \cos(\omega_{\text{mod}} t + \phi)$ and $Q(t) = A_2 \sin(\omega_{\text{mod}} t + \phi)$, we can use the elementary trigonometric identities

$$\begin{aligned} \sin(a) \sin(b) &= (\cos(a - b) - \cos(a + b))/2 \\ \cos(a) \cos(b) &= (\cos(a - b) + \cos(a + b))/2 \end{aligned} \quad (3.19)$$

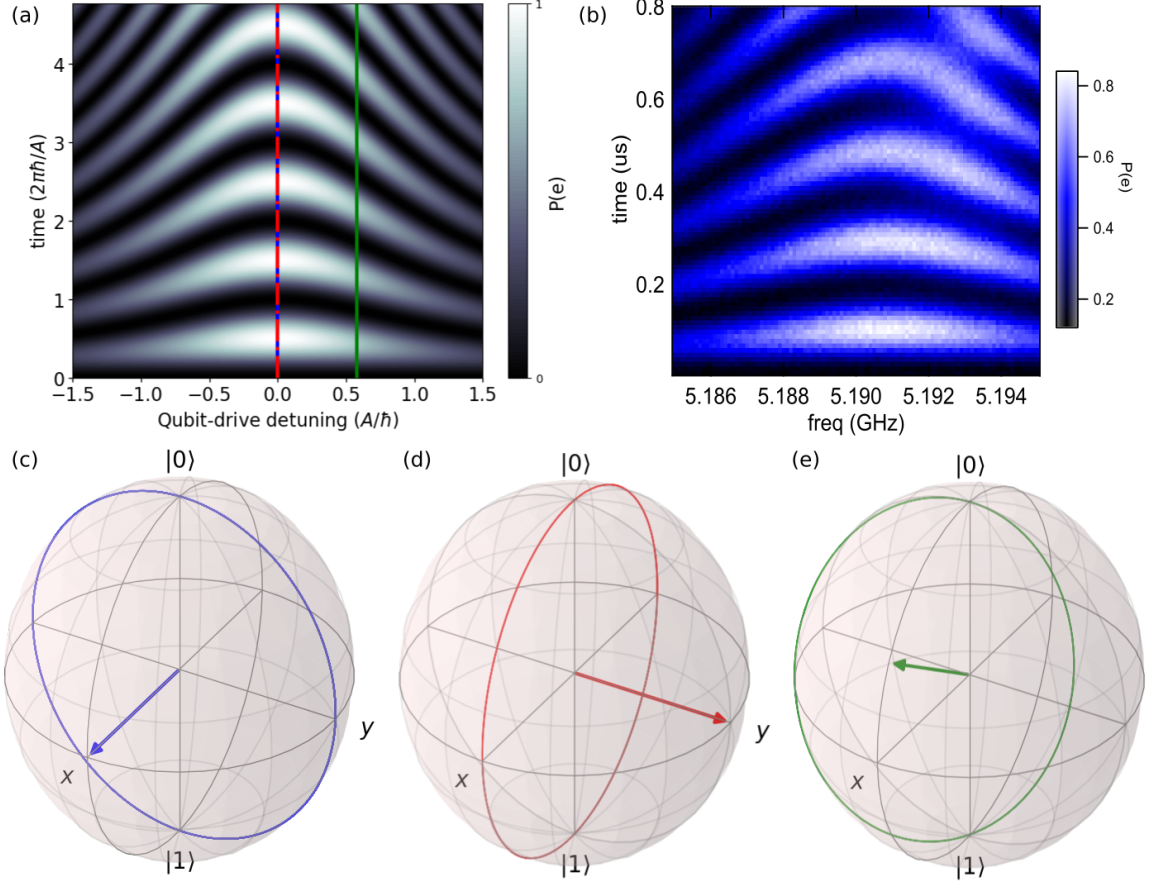


Figure 3.9: **Rabi oscillations:** (a) A plot of the qubit excited state probability as a function of both qubit-drive detuning and drive time, exhibiting the characteristic Rabi “chevron”. (b) An experimental measurement of the Rabi chevron. (c-e) Changing the drive-qubit detuning $\omega_{01} - \omega_d$ and the drive phase ϕ controls the Bloch sphere axis around which Rabi oscillations precess. The quantization axes (arrows) and precession path (lines) are plotted for (c) $\omega_{01} - \omega_d = 0, \phi = 0$ (corresponding to the solid blue line in (a)), (d) $\omega_{01} - \omega_d = 0, \phi = \pi/2$ (dashed red line in (a)), and (e) $\omega_{01} - \omega_d = A/\sqrt{3}, \phi = 0$ (solid green line in (a).) Bloch spheres rendered using the QuTiP package [102].

to show that

$$RF_\phi = A_1 A_2 \cos[(\omega + \omega_{\text{mod}})t + \phi] \quad (3.20)$$

i.e. the phase of the intermediate frequency $I(t)$ and $Q(t)$ signals is imparted onto the high frequency signal out of RF. We can therefore define the axis of Bloch sphere rotation simply by setting the relative phase two sinusoidal waveforms from the Tabor.

To physically implement qubit control pulses, we typically set the continuous wave microwave signal into the LO port to $\omega_{01} - \omega_{\text{mod}}$, and modulate this tone with digitized sinusoidal pulses from the Tabor oscillating at a frequency of $\omega_{\text{mod}} \approx 150$ MHz. One must also “tune” the Tabor to compensate for phase delays in the cables leading to the IQ mixer and imperfections in the IQ mixer, or else tones at the lower sideband ($\omega_{01} - 2\omega_{\text{mod}}$) and LO frequency ($\omega_{01} - \omega_{\text{mod}}$) may leak through the mixer into the experiment. Typically, we adjust the dc offset and relative amplitude/phases of the two Tabor channels leading to the IQ mixer to maximize the upper sideband of the mixed signal (at ω_{01}) and minimize leakage. It is also prudent to remember that mixers are nonlinear elements, and that driving them at excessively high power runs the risk of introducing appreciable signals at $\omega_{01} \pm 2\omega_{\text{mod}}, \omega_{01} \pm 3\omega_{\text{mod}}$, etc... A safe upper bound is ~ 7 dBm into the LO port.

3.4.3 Punchout-based readout

As mentioned in Chapter 2, we may measure the energy eigenstate of the qubit with the dispersive readout technique, monitoring the qubit-state dependent shift in the cavity frequency [60, 61, 103, 68, 34]. While this is entirely possible to do in our lab with only a commercially available high mobility electron transistor (HEMT) amplifier and room temperature amplification, getting high readout fidelity with this method requires optimization of the ratio χ/κ [68, 103, 69]. There turns out to be a more straightforward way to measure the energy eigenstate of the qubit based on the “punchout” mechanism described in Section 3.3.1.

The onset power of the bright state (Fig. 3.6(b)) depends on the state of the qubit [97, 98, 99], meaning that when the qubit is in the $|1\rangle$ state the power required to “punchout” is slightly lower than when the qubit is in the $|0\rangle$ state. We can utilize this fact to make a

high fidelity measurement of the qubit: if we apply a measurement tone to the cavity at the bare cavity frequency ω_q and at a power just below the bright-state threshold when the qubit is in $|0\rangle$, the cavity will “punchout” only when the qubit is in the $|1\rangle$ state. We then map the qubit $|0\rangle$ ($|1\rangle$) states to the low (high) transmission dark (bright) states respectively, which have sufficiently high contrast in the IQ plane to easily be amplified by a HEMT.

In order to map the tone transmitted through the cavity onto the IQ plane, we once again employ an IQ mixer, except this time as a *demodulator*. That is to say IQ mixers also work in reverse: if we feed a tone $A_1 \cos(\omega t)$ into the LO port and a tone $I(t) \cos(\omega t) - Q(t) \sin(\omega t)$ into the RF port, the output of the I port will be $I(t)$ and the output of the Q port will be $Q(t)$. Here, $I(t)$ and $Q(t)$ are slowly varying (with respect to the carrier frequency ω) envelope functions which encode the measured state of the qubit. In this manner, we may reconstruct the position of the microwave pulse on the IQ plane, where the I axis represents the part of the signal in-phase with the LO reference pulse (see Fig. 3.10 for more details.)

Going back to the experimental setup (Fig 3.8), we split a tone at the cavity frequency ω_c into two components. One component feeds into the LO port of the demodulating IQ mixer, where it acts as a reference for the amplified signal coming out of the experiment. The other component feeds into another IQ mixer, which functionally acts as a microwave switch¹¹, allowing signal to pass through when we want to measure the state of the qubit. We then send the output signal through a variable attenuator, which allows us to fine-tune the power of the cavity interrogation tone to just below the “punchout” threshold. The in-line variable phase shifter is optional: we may introduce a phase to make binning the I-Q data in the software somewhat easier.

¹¹For the purposes of the “punchout” readout technique, the IQ mixer here is actually somewhat overkill: we could (and sometimes do) just use a standard microwave switch.

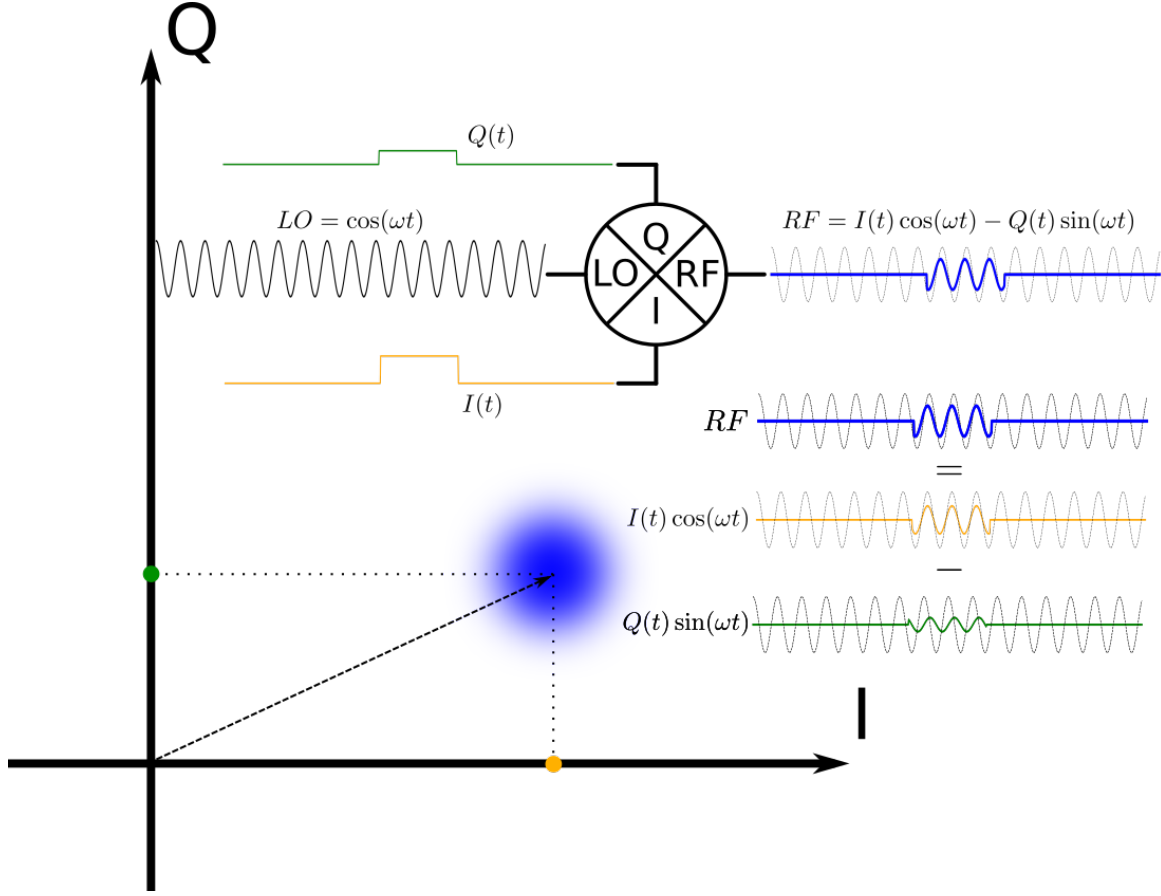


Figure 3.10: **IQ mixer modulation and demodulation:** A microwave tone at the carrier frequency ω is sent into the LO port. If we run the mixer as a demodulator, we send a carrier pulse into the RF port (top, blue) which is composed of in phase (orange, bottom) and out of phase (green, bottom) components relative to the LO signal. If the LO and RF signals are at the same frequency, we record quasi-dc “boxcar” pulses out of the I and Q ports which encode $I(t)$ and $Q(t)$. We use this information to reconstruct the position of the signal on the IQ plane (background.) We may also run the process in reverse, feeding in $I(t)$ and $Q(t)$ to create the desired RF output.

After pulse preparation, the signal is fed into the input line of the experiment, and passes through the input+cavity+output with isolation/amplification as described in § 3.2. The signal is then demodulated using the protocol described above. After one more round of amplification and filtering, the outputs of the I and Q ports are recorded using a two channel analog-to-digital converter (ADC), and the two outputs are plotted in real time as a point on the IQ plane. This constitutes a single measurement of the qubit: after many

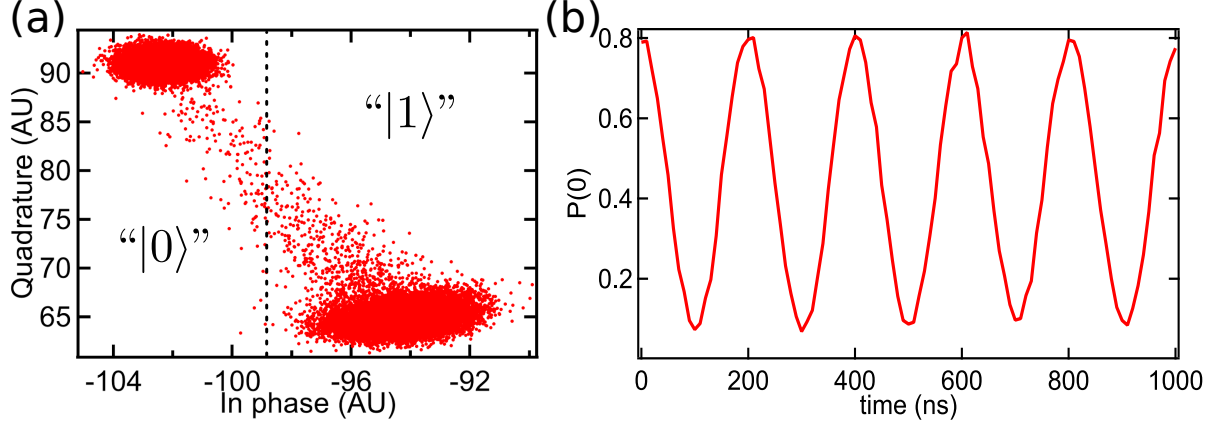


Figure 3.11: **IQ plane measurement and binning:** (a) Example of a series of measurements of the system while Rabi driving the qubit, plotted on the IQ plane. Each point in phase space represents a single punchout measurement of the cavity/qubit system. We bin the measurements by simply defining a cutoff, labeling all points left of the cutoff as $|0\rangle$ measurements and all points right of it $|1\rangle$ measurements. (b) Probability of measuring $|0\rangle$ (point left of the cutoff) as a function of the length of the Rabi drive tone.

measurements, if the experiment is tuned up correctly the points on the IQ plane will form two “blobs”, corresponding to measurements of the qubit in $|0\rangle$ and $|1\rangle$, as in Fig. 3.11(a). We then bin the data and calculate the probability P_0 of the qubit being in the excited state $|0\rangle$ at a given time step. In this manner, we can reconstruct the time dependence of expectation values of the system, allowing us to reconstruct system dynamics/decoherence. An example, is shown in Fig. 3.11(b): here, we apply a variable length tone near the qubit frequency. As the tone length gets longer, the qubit undergoes Rabi flopping from $|0\rangle$ to $|1\rangle$ and back, which we can clearly resolve with a fidelity of $\sim 80\%$.

3.4.4 T_1, T_2 and T_2^e

Having described the experimental setup, we now turn to some of the simplest time-resolved measurements we can do: measuring the energy relaxation and decoherence times of a qubit T_1 and T_2 , as well as the spin-echo time T_2^e . The first step in each of these experiments is to

prepare the measured qubit in the ground state. Since $\hbar\omega_{01} \gg k_B T$, this task is as simple as waiting for time $t \gg T_1$ in between repeated measurements. We rarely measure $T_1 > 30\mu s$, so a repetition rate of several kHz is sufficiently slow to ensure that the qubit has relaxed back to its ground state.

With the qubit in $|0\rangle$, we next apply a π_x -pulse, a microwave control pulse that rotates the Bloch vector by π radian about the x axis, which puts the qubit in the excited state $|1\rangle$, wait for delay time τ , and then measure with a tone at the cavity frequency (Fig. 3.12(a), top.) According to the Bloch-Redfield equations for the evolution of the density matrix (see discussion in § 4.2), we should expect the excited state population to decay exponentially as a function of τ with a characteristic time constant T_1 . In order to measure T_1 , we repeat the experiment varying the delay time τ , and then repeat this sequence of experiments many times. After sufficiently many experiments (typically ≥ 1000), we can calculate the probability that the qubit was found in the excited state as a function of the delay time $P_1(\tau)$, and fit that function to a decaying exponential to extract T_1 (see Fig. 3.12 (a) for an example.)

Measuring T_2 constitutes a similar process: we start off with the qubit in $|0\rangle$, except this time we apply a $\pi_x/2$ pulse, that is to say a microwave pulse that rotates the qubit $\pi/2$ radians about the x axis of the Bloch sphere. This pulse prepares the qubit in $|\psi\rangle = 1/\sqrt{2}(|0\rangle + |1\rangle)$. We then let the qubit evolve for a variable time τ , apply another $\pi_x/2$ pulse and measure. This process is visualized in Fig. 3.12 (b).

The time evolution of a superposition state has two separate components: a deterministic component, which causes the sinusoidal oscillations, and non-deterministic dephasing which causes the amplitude decay. Recall from Eqn. 3.17 that we accomplish a rotation about the x -axis of the Bloch sphere by setting the control pulse phase to $\phi = 0$. This implies we must track the phase of each pulse relative to the first pulse in order to rotate about the

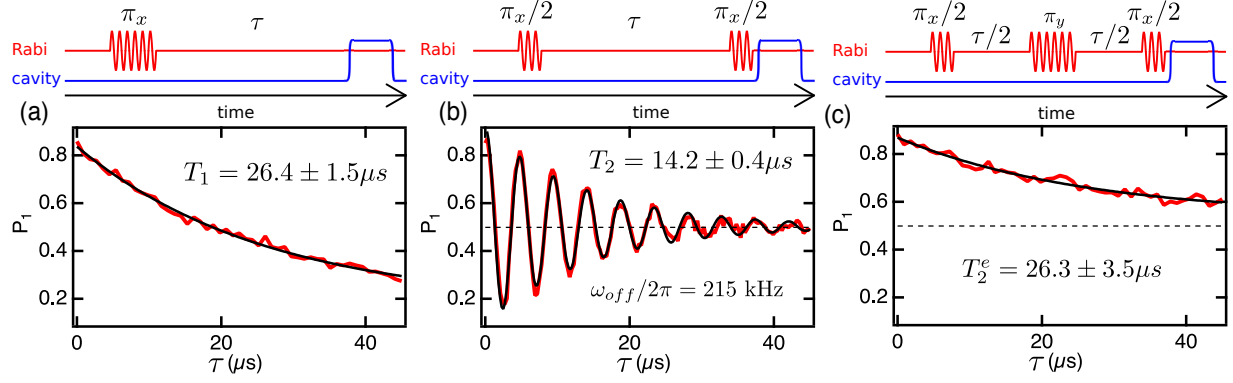


Figure 3.12: **Measuring T_1 , T_2 , and T_2^e :** Pulse sequences and typical data sets for (a) T_1 , (b) T_2 , and (c) T_2^e measurements. In a given measurement, the pulse amplitude/width is fixed, and the waiting time τ is varied between different measurement runs. The sequence is then repeated, and we plot the probability of finding the qubit in the excited state P_1 (red) and fit it to a decaying exponential ((a) and (c), black) or a decaying exponential superimposed on a sinusoid ((b), black) to extract the respective coherence times. We may also extract the qubit-drive detuning $\omega_{off} = |\omega_{01} - \omega_d|$ from the T_2 fit (b).

desired axis of the Bloch sphere: this is another way of staying that *we implicitly work in the frame rotating with the drive pulse*. From Eqn. 3.17 we may read off that when we aren't applying a microwave pulse ($A = 0$), the effective Hamiltonian of the system will be that of a qubit whose transition energy is $\omega_{01} - \omega_d$, meaning the superposition will experience coherent evolution according to

$$|\psi(t)\rangle = \frac{1}{\sqrt{2}}(|0\rangle + e^{-i(\omega_{01} - \omega_d)t} |1\rangle) \quad (3.21)$$

i.e., the state vector will deterministically process around the z -axis of the Bloch sphere at frequency $\omega_{01} - \omega_d$. The action of the second $\pi_x/2$ rotation will depend how much the Bloch vector has rotated: if we apply the second pulse immediately (i.e. $\tau = 0$) it will send the qubit to $|1\rangle$ as if we had just applied a π pulse. If we however wait for the state vector to rotate by π about the z axis (half a period $\pi/(\omega_{01} - \omega_d)$) the qubit will be projected back into $|0\rangle$ instead. Thus, the measured excited state probability P_1 will oscillate sinusoidally

with frequency $(\omega_{01} - \omega_d)/(2\pi)$. Since we know the frequency of the drive signal to high precision, this means that the transition frequency of the qubit ω_{01} is also imparted on the T_2 measurement data.

The qubit will also undergo non-deterministic evolution, which when averaged over many runs will result in the amplitude of the deterministic oscillations decaying exponentially in τ with a characteristic time constant T_2 (again, see § 4.2.) Thus we fit the $P_1(\tau)$ extracted from a T_2 experiment to a decaying exponent superimposed on a sinusoid to extract T_2 and ω_{01} to high precision.

Pure qubit dephasing (at a rate $\Gamma_\phi = \Gamma_2 - \Gamma_1/2$, see § 4.2) may be thought of as stochastic fluctuations in the qubit resonant frequency ω_{01} , caused by elastic (energy conserving) interactions with the environment, which modify the rate at which the qubit precesses around the Bloch sphere [75, 104, 34, 62]. Since these processes are elastic, they are in principle reversible: if we were able to run time “backwards” we would be able to undo the evolution caused by these interactions (as well as any deterministic evolution over the same interval.) We cannot, of course, reverse time. However, if during a T_2 measurement at time $\tau/2$ we apply a π_y pulse (Fig. 3.12(c)), we will invert the position of the Bloch vector about the y -axis of the Bloch sphere sending the relative phase $\phi_{\tau/2}$ accumulated during the first half of the experiment to $-\phi_{\tau/2}$. If the phase change was caused by sources that are slow compared to the experiment time¹², we should expect the accumulated phase in the second half of the

¹²It is reasonable to ask what we mean by “slow with respect to the experiment.” The addition of a π_y pulse can be modeled as adding a *filter function* to the spectral density $S_z[\omega]$ of noise that causes fluctuations in the qubit energy ω_{01} [105, 106, 62]. The filter function, in this case, pushes the frequency range sampled by the experiment up, which is generally beneficial since many sources of noise exhibit a $1/f$ -like spectrum. One may add an arbitrary number of π pulses to the sequence, a technique that was also first developed in NMR called a Carr-Purcell-Meiboom-Gill (CPMG) sequence. In general, the more pulses you add, the higher the frequency range you sample. In fact, it is possible to utilize the fact that the sampled $S_z[\omega]$ depends on pulse number to work backwards and reconstruct $S_z[\omega]$ using a series of CPMG measurements with different numbers of pulses, as was done in [106].

to *also* change by $\phi_\tau/2$, rendering the total phase $\phi_\tau = 0$. This phase refocusing experiment, known as a “spin-echo” experiment in nuclear magnetic resonance (NMR) where it was first developed [107], “undoes” deterministic phase evolution and the phase evolution caused by slowly varying noise sources, which removes the sinusoidal pattern in T_2 data and may extend the dephasing time. This type of experiment is so common that the decay constant observed during it gets its own name: T_2^e (which is shorthand for T_2 -echo.) T_2^e should be $\geq T_2$, however since we cannot undo spontaneous decay by refocusing the Bloch vector, we are still limited by $T_2^e \leq 2T_1$.

Chapter 4

Integrating superfluids with superconducting qubit systems

When we built a description of quantum circuits in Chapter 2, we used the fact that the superconducting gap shielded the qubit from single-particle excitations to sweep the microscopics under the rug and write down phenomenological Hamiltonians in terms of macroscopic circuit parameters like the E_C and E_J . We also required that the system be cooled to dilution refrigerator temperatures, where $k_b T \ll \hbar \omega_{01}$, to exponentially suppress any thermal excitations to the system. Then, in Chapter 3, we dedicated quite a bit of time to discussing strategies to mitigate decohering noise and spurious excitations that limit our ability to perform complex control operations on the system. Ideally these steps would be enough to protect superconducting qubits from decohering noise.

In reality, however, the microscopic details of the Hamiltonian become important for decoherence. Microscopic fluctuations may cause fluctuation in the circuit parameters, which may in turn cause fluctuations in ω_{01} and dephase the qubit. Microscopic quantum systems with transition energies near $\hbar \omega_{01}$ may provide a density of states for the qubit to decay into, and these systems will also fluctuate in unpredictable ways. Additionally, various degrees of freedom that couple to the qubit may be poorly thermalized to the cryostat, with background excited state occupancies much larger than naively expected from the

mixing chamber temperature $T_{MXC} \approx 10$ mK. Indeed, the MXC thermometer is at best a proxy for the temperature of *electrons in the copper MXC plate*: it cannot be assumed that the temperature of a given thermal bath mechanically connected to the MXC is at this temperature.

From this perspective, it would be advantageous to have some way of *locally* thermalizing various heat baths that couple to the system. Looking to other areas of low temperature physics, liquid helium is commonly employed as a passive thermalizing agent for experiments requiring extremely low temperatures. For example, helium immersion cells are used in the study of two-dimensional electron systems in the quantum Hall regime to more efficiently thermalize these systems at milli-Kelvin temperatures [108, 109]. However similar techniques have not been employed in the setting of superconducting circuits.

There is also growing fundamental interest in studying the mechanical motion of superfluid helium at the quantum limit. Recent experiments and proposals have investigated superfluid helium as a platform for optomechanical experiments [110, 111, 112, 113], or as a substrate for an electron motional qubit [114, 115, 116, 117, 118, 119], where details of the superfluid surface mechanics are important to understanding the decoherence of the proposed qubit. More in line with the main topic of this dissertation, one may also ask whether it is possible to *coherently* control mechanical excitations of the superfluid at the single quantum level, similar to what has been achieved in solid-state mechanical resonators [2, 10, 5, 4]. When considering both the study and control of mechanical excitations at the single quantum level, the flexibility of superconducting circuits offers a distinct advantage to engineering systems that preferentially couple to specific modes.

With both these motivations in mind, this chapter details an experiment systematically studying the effects of superfluid helium on a “standard” 3D-transmon superconducting

qubit/cavity system. The main focus of the experimental results is understanding how the superfluid helium influences the spectral and decoherence properties of the cavity/qubit system. At the end of the chapter, we will loop back around and discuss how one may or may not be able to use superconducting circuits to study/manipulate mechanical excitations in quantum fluids. The results in this chapter were also reported in Ref. [120].

4.1 Superfluid helium

The story of superfluid helium once again begins with H. Kamerlingh Onnes, who in 1908 became the first person to liquify the noble gas helium¹ by cooling it to 4.2 K. In fact, the technical accomplishment of liquifying helium was what allowed Kamerlingh Onnes to study the resistance of low-temperature metals, leading to the discovery of superconductivity: to this day, liquid helium is indispensable as a cryogenic refrigerant. Every method for cooling bulk matter below 10 K uses helium in some fashion [89]. Since most conventional superconductors also require these temperatures to operate, industrial or scientific processes that use superconductors, from MRI machines to particle accelerators, are also completely dependent on refrigeration from liquid helium.

From a physics perspective, however, liquid helium has far more interesting properties than a low evaporation temperature. For starters, at atmospheric pressure liquid helium never freezes: it remains a liquid *all the way down to absolute zero*. This astonishing fact may be understood in terms of zero point fluctuations in the motion of helium atoms.

¹Helium has two stable isotopes: ^4He , which is largely formed in Earth's crust by alpha decay of heavy radioactive elements and extracted from natural gas, and much less abundant ^3He , which is usually produced artificially by tritium decay. These isotopes have very different low temperature properties, owing to their significant mass difference and also the fact that ^3He is a fermion while ^4He is a boson. However, ^3He only exists in trace quantities in natural helium reserves, and therefore we will focus exclusively on ^4He in this chapter, unless noted otherwise.

Helium atoms are very light, and the van der Waals interaction between two helium atoms is very weak. In a hypothetical helium solid, an individual helium atom would be held in place by a 3D harmonic oscillator potential. However, because of the light mass and weak intra-particle interaction, the zero point energy fluctuations $E_{ZPF} = 3\hbar\omega/2 \propto m^{-1/2}$ would far outstrip the attractive van der Waals potential, rendering such a solid unstable [121]. Even at the lowest temperatures, helium does not form a solid unless pressurized above ~ 25 atmospheres!

However, upon cooling to low enough temperatures, helium does not remain a typical fluid. In the late 1920's, it was discovered that the specific heat of liquid ^4He exhibits a dramatic and narrow peak as a function of temperature around $T_\lambda = 2.17$ K, indicating the existence of a phase transition at this temperature. The following decade saw a series of experiments building evidence that, below T_λ , liquid helium does not behave like a normal fluid, but rather exhibits many strange phenomena, especially when flowing through tight constrictions such as a thin capillary tube or the microscopic channels in a tightly packed powder [122]. These strange observations were eventually reconciled by modeling liquid helium below T_λ as a sum of two component fluids: a normal fluid with density ρ_n and a *superfluid* with density ρ_s , obeying

$$\rho_{He} = \rho_n + \rho_s \tag{4.1}$$

where ρ_{He} is the total density of liquid helium. The superfluid component flows with precisely zero viscosity, and carries no entropy, while the normal component behaves more or less like a normal fluid (i.e. with finite viscosity and entropy.) This two-fluid model [123] explains the behavior of the fluid flow through tight constrictions: the non-viscous superfluid may

easily flow through such a “superleak”, while the viscous normal component is impeded so much that the flow is effectively zero.

Superfluidity turns out to be closely related to Bose-Einstein condensation, where a macroscopic number of atoms occupy the ground state of the system. The theoretical description of superfluid helium is, however, more complicated than that of Bose-Einstein condensation owing to the strong intra-particle interaction of the densely packed helium atoms [121]. For the purposes of this dissertation, however, it is sufficient to note that well below the transition temperature T_λ , the normal fluid density is empirically found to be [121]

$$\rho_n \propto T^4 \tag{4.2}$$

Thus, at the temperatures relevant to superconducting circuits, $\rho_n \approx 0$ and $\rho_s \approx \rho_{He}$. For the experimental portion of this chapter, we will also ignore the direct exchange of excitations of the qubit with mechanical modes of the helium. In this approximation, the superfluid becomes, for all intents and purposes, an extremely low loss dielectric² that may conduct heat away from various thermal baths coupled to the cavity/qubit system [89, 125].

4.2 Decoherence and noise

Much of the experimental result below focuses on the decoherence properties of the qubit, and how the presence of superfluid helium modifies this decoherence. Strictly speaking, decoherence is the process by which the system we care about (the qubit) interacts with

²High power measurements in the context of accelerator physics have yielded an upper bound on the loss tangent of liquid helium $\tan \delta = \text{Im}\{\epsilon\}/\text{Re}\{\epsilon\} < 10^{-10}$ at ~ 1.5 K [124].

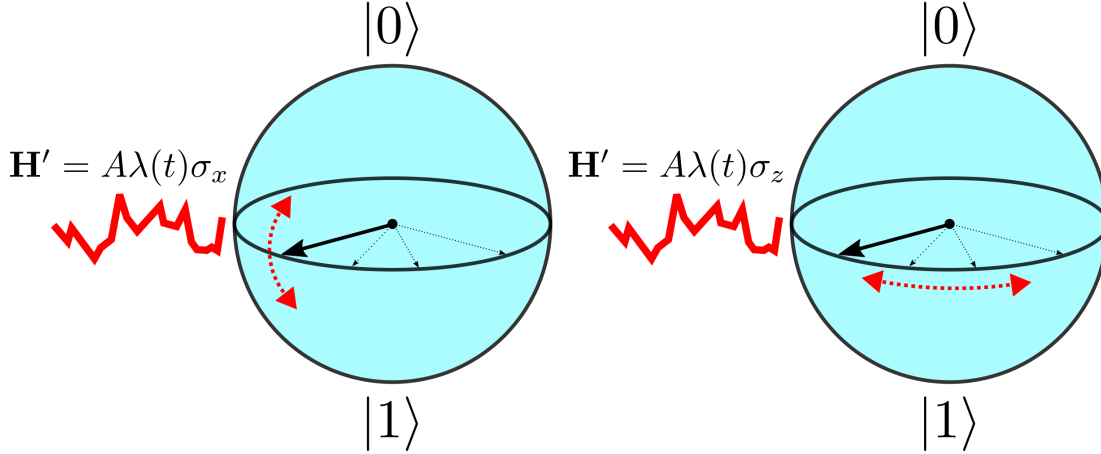


Figure 4.1: **Noise and decoherence of a qubit:** Left: depolarizing noise (sometimes called transverse noise) of a qubit may be thought of as spurious environmental interactions that kick the Bloch vector (red dashed arrow) in a direction transverse to the deterministic evolution of the qubit (black dashed arrows.) This noise is introduced by terms in the qubit Hamiltonian $\propto \sigma_x$ or σ_y . Right: Dephasing noise (sometimes called longitudinal noise) of a qubit may be thought of as spurious environmental interactions that kick the Bloch vector along the equator, causing shifts relative to the deterministic phase evolution of the qubit. This noise is introduced by terms in the qubit Hamiltonian $\propto \sigma_z$.

and becomes entangled with its environment, a many-DOF system that we have no hope of tracking. From this perspective, we can think of entanglement as information *loss* into the environment: the fluctuations in the environment take a well characterized quantum state and, over some timescale, kick it around such that the information we once had about the state is now useless. Mathematically, this is modeled as a *pure* state (a state that may be written down as a wave function) undergoing *non-unitary* evolution and decaying into a *mixed* state (a state which we're forced to write as a density matrix, since we have incomplete knowledge of it.)

The process of non-unitary evolution is often modeled by a “master” equation, or Lindbladian, however the decay of a quantum two-level system may often be described by the Bloch-Redfield approach [126, 127]. In the Bloch-Redfield picture, the non-unitary evolution

(i.e. decoherence) of a quantum two-level system characterized by two rates, the **depolarization rate** Γ_1 and the **dephasing rate** Γ_2 . The depolarization rate, often quoted as the reciprocal of the depolarization time $T_1 = \Gamma_1^{-1}$, describes how quickly we lose information about the relative magnitudes of a $|0\rangle$ and $|1\rangle$ in a superposition state. Information loss about the magnitudes comes about because of *inelastic* scattering, i.e. the qubit exchanging excitations with its environment. Since rates add, Γ_1 is given by

$$\Gamma_1 = \Gamma_{\uparrow} + \Gamma_{\downarrow} \quad (4.3)$$

where $\Gamma_{\uparrow}$ is the rate at which the qubit is spuriously excited by the environment and $\Gamma_{\downarrow}$ is the rate at which the qubit in the excited state loses an excitation to its environment. In the limit $\hbar\omega \gg k_B T$, spurious excitations are exponentially suppressed ($\Gamma_{\uparrow} \ll \Gamma_{\downarrow}$) and Γ_1 is given by Fermi's golden rule³ [104, 34, 93, 62]:

$$\Gamma_1 \approx \Gamma_{\downarrow} = \frac{2\pi}{\hbar} \left| \langle 0 | \hat{H}' | 1 \rangle \right|^2 \rho(\hbar\omega_{01}) \quad (4.4)$$

where here, $\rho(\hbar\omega_{01})$ is the density of states at the qubit transition energy and $\hat{H}'$ is a perturbative term in the Hamiltonian induced by some noise source. The energy decay rate is also called the *longitudinal* decay rate, since the only perturbative terms that cause relaxation are terms $\propto \sigma_x, \sigma_y$, i.e. terms that rotate the Bloch vector longitudinally relative to the quantization axis.

The dephasing rate Γ_2 quantifies how quickly we lose information about the relative phases of a $|0\rangle$ and $|1\rangle$ superposition state. Recall the argument we made in Chapter 2

³Fermi's golden rule may be equivalently formulated in terms of a fluctuating quantity λ that couples to the qubit via a perturbative Hamiltonian term $\mathbf{H}' = A\lambda\sigma_x$. In this formulation $\Gamma_{\downarrow,\uparrow} = (A/\hbar)^2 S_{\lambda\lambda}[\pm\omega_{01}]$, where $S_{\lambda\lambda}[\omega]$ is the autocorrelation function of the noisy prefactor λ at frequency ω [93].

when discussing the original justification for the transmon qubit: an equal superposition state evolves as a function of time $|\psi(t)\rangle = 1/\sqrt{2}(|0\rangle + e^{-i\omega_{01}t}|1\rangle)$. If ω_{01} is constant, we can account for this dynamic phase by working in the frame rotating with the qubit, however if ω_{01} fluctuates as a function of time, the phase will evolve non-deterministically, and we will lose information about the system. We may *also* lose phase information about the system from depolarizing processes: the dephasing rate Γ_2 is the sum of a contribution from depolarization and a “pure dephasing” rate Γ_ϕ quantifying fluctuations in ω_{01} :

$$\Gamma_2 = \frac{1}{2}\Gamma_1 + \Gamma_\phi \quad (4.5)$$

Note that even with little pure dephasing ($\Gamma_\phi \rightarrow 0$), we are still limited by $\Gamma_2 \geq \Gamma_1/2$ (or, equivalently, $T_2 \leq 2T_1$.) Noise that causes changes in ω_{01} effectively adds a perturbation to the qubit Hamiltonian $\propto \sigma_z$. Therefore, Γ_2 is also called the *transverse* decay rate, since terms $\propto \sigma_z$ rotate the Bloch vector along its transverse (equatorial) axis.

When $\hbar\omega_{01} \gg k_bT$ ($\Gamma_1 \approx \Gamma_\downarrow$), the Bloch-Redfield approach predicts that the density matrix ρ of a qubit with the initial state $|\psi(t=0)\rangle = \alpha|0\rangle + \beta|1\rangle$ will evolve as [104, 62]

$$\rho = \begin{pmatrix} 1 + (|\alpha|^2 - 1)e^{-t/T_1} & \alpha\beta^*e^{i\omega_{01}t-t/T_2} \\ \alpha^*\beta e^{-i\omega_{01}t-t/T_2} & |\beta|^2e^{-t/T_1} \end{pmatrix} \quad (4.6)$$

i.e. the excited state probability ρ_{11} will decay exponentially with a time constant T_1 , while the off-diagonal terms will decay exponentially⁴ with a time constant T_2 .

⁴This simple exponential decay is conditional on the noise source driving decoherence being weak enough to be treated perturbatively and not having long timescale correlations. This second condition is violated for common $1/f$ noise, where $S_{\lambda\lambda}[\omega] \propto 1/\omega^n$. Violation of this assumption can result in a non-exponential decay envelope [104], however all the data in this chapter fit a decaying exponential reasonably well and we will ignore these subtleties.

4.2.1 Common noise sources in superconducting qubits

To facilitate some of the discussion below, we offer a brief survey of noise sources that commonly drive decoherence in superconducting circuits.

Charge noise: In Chapter 2, we briefly discussed charge noise, i.e. the uncontrolled fluctuation of the microscopic electrostatic environment of the qubit. These fluctuations may be caused a number of sources: shot noise in nearby conductors, fluctuations in charged lattice defects in the substrate, fluctuations in adsorbed surface molecules, etc. However, since transmon qubits (and most modern superconducting qubits for that matter) are exponentially insensitive to charge noise, we will ignore charge noise for the present discussion.

Flux noise: In the case of the flux tunable transmon (or other flux tunable qubits, such as flux qubits or fluxonium), fluctuations in the magnetic flux Φ_{ext} threading the SQUID loop may cause cause fluctuations in ω_{01} , driving dephasing. Flux noise is thought to be caused by both macroscopic sources (such as current fluctuations in the flux-tuning coil) and microscopic sources (such as spins at the metal-insulator surface [128].) In the experimental data presented below, we use a fixed frequency transmon, and thus flux noise is of little concern for us.

Dielectric loss: Like all classical integrated circuits, superconducting circuits must be fabricated on a dielectric substrate. All real dielectric materials absorb some fraction of an impinging electric field: this absorption may be modeled by making dielectric constant ϵ imaginary, where the $\text{Im}\{\epsilon\}$ quantifies the electrical loss of the material. A common figure of merit of a dielectric is the loss tangent: $\tan(\delta) = \text{Im}\{\epsilon\}/\text{Re}\{\epsilon\}$. The two most common substrates superconducting circuits are fabricated on, high resistivity silicon and sapphire, have very low $\tan(\delta) < 10^{-6}$ [79], though there is evidence that sapphire is the better of the

two [129]. In addition to the host substrate, interfaces (say, between metal and substrate or metal and air) are thought to host thin, amorphous layers of oxides and organic materials that contribute significantly to loss [130, 79, 129, 131].

On the quantum level, dielectric loss is likely caused by a large ensemble of defects and impurities that effectively create a bath of quantum two-level systems (TLS) [132]. Though any one of these TLSs is weakly coupled to the qubit mode, they effectively create a density of states for the qubit to decay into. These ensemble losses are, however, largely a function of the fabrication details. Since our experiment provides a comparison of one device subject to systematically different conditions, we will assume these details are constant over the data presented below.

Two-level systems: In a qualitatively different regime, individual TLSs near or in the Josephson junction may strongly couple to the qubit degree of freedom, driving decoherence [133]. If a TLS is *inside* the Josephson tunnel barrier, a shift in the state of the TLS may locally change the microscopic details of the tunneling barrier, slightly modifying the critical current and thus E_J . These TLSs are thought to cause dephasing [134], and have been blamed for discrete shifts in the qubit frequency that persist over long timescales [40].

Additionally, strongly coupled TLSs with transition frequencies near ω_{01} provide a density of states for the qubit to decay into. It has recently been shown that these TLSs, which themselves lack the transmon's protection from quasistatic charge noise, undergo random walks in frequency space, moving in and out of resonance with the qubit on $\sim$ hour timescales [135]. As the TLS moves in and out of resonance, an environmental interaction channel is effectively turned on and off, which is thought to cause long timescale changes in the qubit coherence properties that pose a problem for benchmarking many-qubit systems [136, 137, 138]. These fluctuations may be impacted by the presence of thermalizing superfluid helium

[125], and strongly coupled TLSs are of interest to us in this study.

Photon shot noise: In Chapter 2, we saw that when the cavity and qubit are dispersively coupled, the effective resonant frequency of the cavity depends on the state of the qubit. Conversely, we could take Eqn. 2.46 and regroup terms:

$$\mathbf{H}_{\text{JC,disp}} \approx \hbar\omega_c(\mathbf{a}^\dagger\mathbf{a} + \frac{1}{2}) + \frac{\hbar}{2}\left[\omega_q + \frac{g^2}{\Delta}\mathbf{a}^\dagger\mathbf{a}\right]\boldsymbol{\sigma}_z \quad (4.7)$$

From this perspective, it looks like the effective frequency of the *qubit* depends on the cavity photon number occupancy [68, 37]. The dispersive shift, written in this way, implies that if the cavity photon population is fluctuating (i.e., if the cavity mode is poorly thermalized), these fluctuations will also drive dephasing in the qubit [76, 87, 88, 90].

Quasiparticles: In § 2.2.4, we placed a great deal of emphasis in working in conditions where electronic excitations above the superconducting ground state (referred to as superconducting quasiparticles, or simply “quasiparticles” in the literature) are suppressed. The expected density of quasiparticle excitations should be $\propto e^{-\Delta_{\text{BCS}}/k_B T}$: for aluminum at 10 mK, this factor is vanishingly small. However, many experiments [139, 140, 141, 142] have recorded quasiparticle densities many orders of magnitude higher than the naive prediction from BCS theory, such that loss from quasiparticles (sometimes referred to as “quasiparticle poisoning”) is a significant limiting factor in qubit coherence.

The exact origin of these elevated quasiparticle levels is still subject to some debate. It is agreed upon that improper filtering may introduce pair-breaking photons from warmer parts of the cryostat, and we discussed several ways to mitigate these effects in Chapter 3. However, even with highly efficient filtering, there is a growing body of evidence that nonequilibrium quasiparticles may be generated in significant quantities by decay of radioactive isotopes in

the experimental setup [143], cosmic rays impinging on the substrate [95, 144], or potentially “heat leaks” in the cryostat that take a very long time ($\sim$ days-months) to thermalize [96]. Quasiparticles may also be excited by/convert into high frequency phonons in the underlying substrate [145], suppressing recombination and lengthening the lifetime of deleterious excitations. Suppressing quasiparticle sources, and figuring out how to thermalize these excitations more efficiently, is an open area of research in the field.

4.3 Experimental setup

4.3.1 Gas handling system and fill line

Introducing a controllable level of superfluid helium into an experiment introduces several experimental complications that must be solved. First and foremost, since superfluid helium has zero viscosity, it will leak through more or less any porous structure, and through many crevices that normal fluid would not leak through. Identifying and preventing superleaks is an age-old pastime of low temperature physicists, and the experimental setup must be carefully designed to minimize the risk of superleaks developing.

Additionally, one must consider the thermalization of superfluid helium, normal liquid helium, and helium vapor at the various stages of the cryostat. Care must be taken to ensure that the helium thermalizes to the base temperature of the cryostat (~ 10 mK.) On the other hand, the superfluid film flow will cause a thin superfluid layer to form on the inside of any surface exposed to helium below the superfluid transition temperature. Since superfluid helium is a good thermal conductor, the film flow effect puts us in danger of thermally connecting the lower stages of the cryostat, which may result in unwanted heat flow into the MXC.

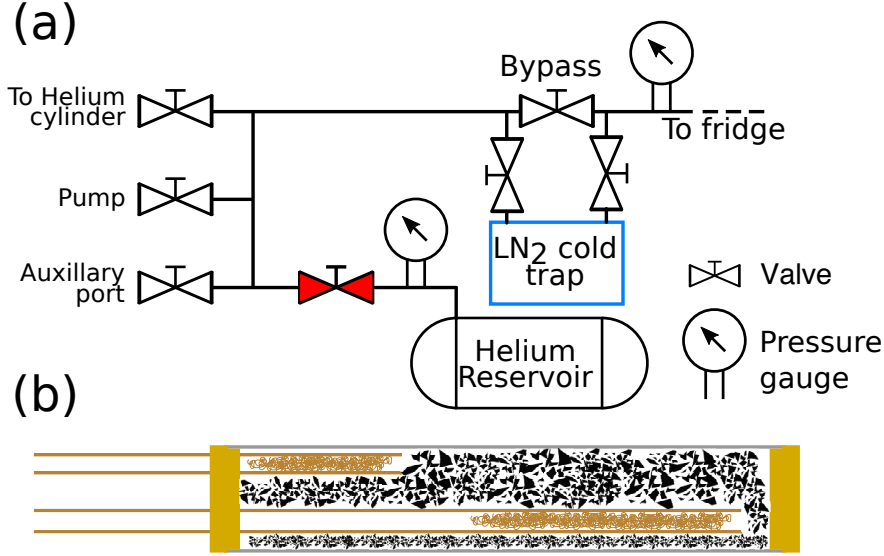


Figure 4.2: **Gas handling system and cold trap:** (a) Schematic of the helium gas handling system. The red valve represents a compound shutoff-throttling valve to controllably introduce helium into the sample cell. All helium introduced into the cell should be passed through the cold trap. (b) Cut-away of the cold trap. Two copper tubes are fed into a stainless steel cylinder with brass caps hard soldered on the end. The cold trap is filled with activated charcoal, which has a high surface area that adsorbs atmospheric gasses when chilled to liquid nitrogen temperatures. The copper tubes are filled with a copper mesh to act as another adsorbing surface and to prevent charcoal debris from exiting the cold trap.

We must also be extremely careful to prevent the introduction of non-helium gas into the fill line. Even at the 50 K stage of the cryostat, all common atmospheric gasses will have solidified: if there is a leak to atmosphere, one runs the danger of blocking the fill line and leaving no escape route for the liquid helium already in the fridge, which could lead to extreme pressure buildups.

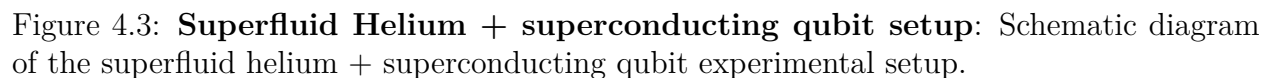
In Figure 4.2 (a) we present a schematic diagram of the gas handling system used to controllably and safely introduce helium into the experiment. We introduce helium into the experiment from a helium reservoir with a known volume through the throttling valve (red, Pfeiffer EVN 116), which precisely controls the flow rate. The pressure gauge on the helium reservoir gives us a rough estimate of how much helium has been introduced into the fridge.

When filling the experimental volume with helium, one should always fill through the liquid nitrogen cold trap (Fig 4.2 (b)), which provides a large, cold surface area for atmospheric gasses that may have leaked into the gas handling system to adsorb onto. It is also best practice keep the cold trap cold for the duration of the experiment and open to both the fill line and the (emptied) reservoir: this will hopefully prevent atmospheric gas that leaks into the system from blocking the fill line, and in the event of a pulse tube shutoff provides a safe path to the reservoir for the helium to expand into when it boils and exits the cell.

After filtering through the liquid nitrogen cold trap, the helium enters the fridge through a series of hermetically sealed tubes that constitute the helium fill line. The fill line must also thermalize the helium at each stage of the dilution refrigerator in transit to the sample cell, a hermetically sealed 3D microwave cavity (see Fig. 4.3 for a schematic of the fill line.) Inside of the fridge, the capillary line is a 1/16" diameter stainless steel capillary everywhere with the exception of two sections: the section between room temperature and the 50 K plate (made of 1/8" brass tube⁵), and the section between the cold plate and mixing chamber plate, where a 0.017" CuNi capillary is used to minimize heat flow between these two plates by a superfluid film within the capillary. The incoming helium is thermalized at five points: twice by mechanically clamping the fill capillary to the 50 K and 4 K plates of the cryostat, and by passing the helium through a copper sinter heat exchanger at each of the still plate (800 mK), cold plate (100 mK), and mixing chamber plate (MXC). In this configuration, we find that the lowest temperature of the dilution refrigerator is not substantially changed upon filling the 3D cavity with superfluid ^4He .

For ease of installation and repair, the fill line is broken up at several points in the fridge. The first connection inside the vacuum can, made between the 1/4" copper tube feeding

⁵We also put copper mesh in the end of the brass section to make a secondary cold trap



into the fridge and the 1/8" brass tube, is a brass Swagelok joint. Every other joint in the fridge is either a soft solder joint or a joint consisting of two brass flanges sealed with an indium o-ring (see Fig. 4.3). It is imperative to thoroughly leak check each joint with a helium leak-checker upon first making the connection. If there is a superleak inside of the fridge during a cool down, superfluid helium that has escaped from the fill line will flow to the warmest point in the fridge, evaporate and thermally short every stage of the fridge to room temperature and causing the fridge to warm up *very* quickly. This is, needless to say, a scenario worth avoiding. Note that the flow impedance through the sinter heat exchangers is large, rendering it difficult to leak check the parts of the line that come after them. It is therefore advisable to leak check the sample cell before one puts it on the fridge.

4.3.2 Superfluid helium leak tight sample cell

The experiment consists of a single-junction transmon circuit, fabricated by Kater Murch's group at Washington University in St. Louis, housed in a rectangular 3D aluminum microwave cavity (see Fig. 4.4(a).) The two halves of the cavity are hermetically sealed with a conventional indium wire o-ring typically used for making superfluid leak tight joints. The external microwave coupling to the cavity is provided via two hermetically sealed $50\ \Omega$ assemblies (see Fig. 4.4(b).) These assemblies consist of commercial hermetic GPO feedthroughs⁶ in which the room temperature rubber o-ring has been replaced with a cryogenic indium seal [146, 147]. These GPO feedthroughs are seated in custom brass flanges that are themselves sealed to the body of the cavity with indium o-rings.

Coupling pins (Fig. 4.4(b)(5)) are soldered to the inner portion of each assembly to provide microwave signals to the cavity. The coupling pin is made out of 16 gauge copper

⁶Gilbert Engineering part # 0119-783-1

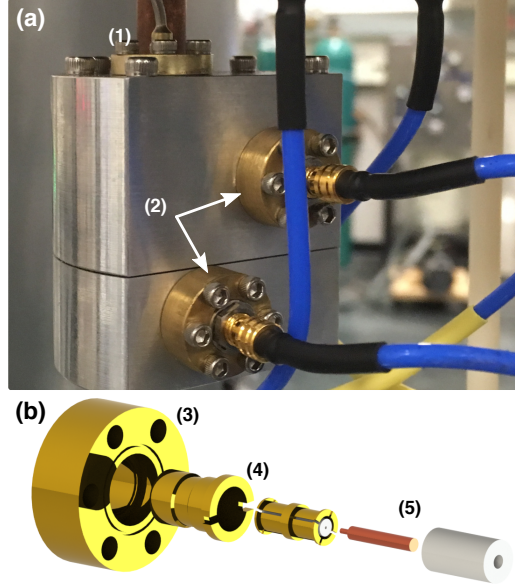


Figure 4.4: **Superfluid leak-tight sample cell, iteration #1:** (a) Picture of the hermetically sealed 3D superconducting microwave cavity. Visible are (1) the helium fill capillary and flange and (2) the two microwave coupling ports. (b) Exploded rendering of the custom microwave coupling assembly. The hermetic GPO feedthrough (4) sits in a brass flange (3), and is sealed with an indium o-ring in between both the feedthrough and the flange, and the flange and the wall of the 3D cavity. A $50\ \Omega$ impedance matched copper pin (5) is soldered into a GPO “bullet” connector and extends to the inner wall of the cavity and provides coupling to the TE₁₀₁ fundamental mode of the cavity.

wire, which is substantially thicker than the inner pin of a standard GPO connector. To fit the pin into the GPO connector, one end was etched in ferric chloride for ~ 20 minutes to decrease its diameter. A small Teflon cylinder fits around the port to match⁷ the pin to $50\ \Omega$. The end of the pin on the input port is flush with the cavity wall, while the pin on the output port protrudes into the cavity to increase coupling to the output. The microwave ports are connected to a standard filter/amplifier chain (see Fig 4.3.) The cavity is thermally anchored to the mixing chamber plate of the dilution refrigerator, and helium is introduced via a hole on the top of the cavity that connects to a stainless steel fill capillary through a custom brass flange, which is itself hermetically sealed to the cell with an indium o-ring.

⁷This was likely superfluous, since the cavity itself is not impedance matched to $50\ \Omega$. More recent experiments have dropped the Teflon.

4.4 Experimental results

4.4.1 Cavity and qubit spectroscopy

To characterize the effect of adding superfluid ^4He we first perform continuous wave spectroscopy of the cavity/qubit coupled system, both when the cavity is empty and under vacuum, and when it is filled with superfluid helium. Using a vector network analyzer we characterize the cavity response by measuring the microwave transmission (S_{21}) through the measurement circuit as a function of frequency. At high power (~ -80 dBm power injected into the cavity), the measured response is Lorentzian and peaked at the classical cavity fundamental frequency [98, 97, 99] $f_c = \omega_c/2\pi$, shown as the blue (dark gray) traces in Fig. 4.5(a). The change in the speed of light caused by the presence of a dielectric of relatively permittivity ϵ should shift the bare cavity frequency from $f_c \rightarrow f_c/\sqrt{\epsilon}$. Indeed, we find that, when helium is added to the cavity, the fundamental frequency f_c shifts from 6.93480 GHz to 6.75395 GHz (see Table 4.1), corresponding to an effective cavity dielectric constant of $\epsilon = 1.054$, which agrees well with that of superfluid helium $\epsilon_{He} = 1.057$ [148, 149]. We also note that the quality factor of the microwave resonator is not significantly affected by the presence of helium, consistent with the findings in Refs. [150, 110, 148].

The hybrid cavity/qubit system may be described by the generalized Jaynes-Cummings Hamiltonian (JCH), which takes into account the higher excited states of the transmon circuit $|i\rangle$:

$$\hat{H}_{JC} = \hbar\omega_c \mathbf{a}^\dagger \mathbf{a} + \sum_i \omega_i |i\rangle \langle i| + \hbar \sum_i \left(g_{i,i+1} |i\rangle \langle i+1| \mathbf{a}^\dagger + \text{h.c.} \right). \quad (4.8)$$

In Eq. 1 $\mathbf{a}^\dagger$ and $\mathbf{a}$ correspond to the microwave cavity photon creation and annihilation

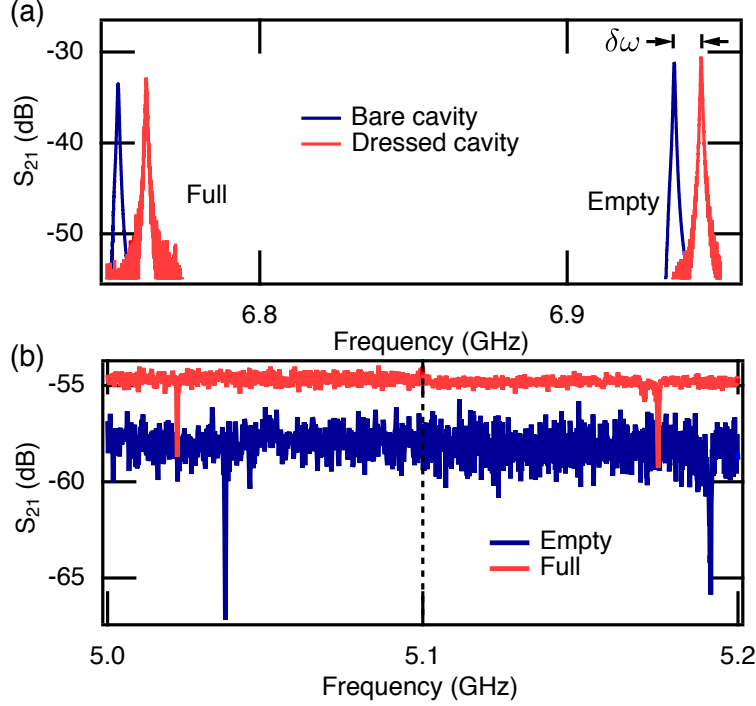


Figure 4.5: **Cavity and qubit spectroscopy is the presence of superfluid helium:** (a) Measured cavity transmission S_{21} as a function of frequency when the cavity is empty (right) and full of superfluid helium (left). Depending on the level of the applied microwave power we can measure both the cavity resonance dressed by the qubit in its ground state (red/light gray, $P \approx -120$ dBm) or the bare cavity resonance (blue/dark gray, $P \approx -80$ dBm.) (b) Two-tone spectroscopy of the qubit immersed in liquid helium (red/light gray) and in vacuum (blue/dark gray), offset vertically for clarity. Right of the dotted line, a low power tone is applied to excite the qubit ($P_q \approx -120$ dBm) from its ground state $|0\rangle$ to its first excited state $|1\rangle$, while to the left of the line a high power ($P_q \approx -90$ dBm) tone is applied to induce a two photon transition from $|0\rangle$ to $|2\rangle$. The dips in the transmission correspond to qubit excitation frequencies ω_{01} (right) and $(\omega_{01} + \omega_{12})/2$ (left).

operators respectively, and h.c. stands for hermitian conjugate. In the transmon regime [34], the uncoupled qubit frequencies ω_i are determined by the Josephson energy E_J , the charging energy E_C , and the cavity/qubit couplings constants $g_{i,i+1} \approx g_{01}\sqrt{i+1}$. In this limit, Eqn 4.8 is therefore determined by E_J , E_C , the ground-to-first excited state vacuum Rabi splitting g_{01} , and the cavity frequency ω_c .

In addition to shifting the cavity resonant frequency, the presence of dielectric superfluid will also modify all of the spectroscopic parameters of the coupled qubit/cavity system,

Value	Empty (GHz)	Full (GHz)	change (%)
$\omega_c/2\pi$	6.9348	6.7540	-2.62
$\delta\omega/2\pi$	0.00875	0.00913	4.32
$\omega_{01}/2\pi$	5.1914	5.1747	-0.32
$\omega_{12}/2\pi$	4.8834	4.8695	-0.28
E_J/h	13.887	13.895	0.06
E_C/h	0.2710	0.2690	-0.82
$g_{01}/2\pi$	0.1235	0.1201	-2.8

Table 4.1: Spectroscopic parameters of the cavity/qubit system both in the presence and absence of superfluid helium. $\omega_c, \delta\omega, \omega_{01}$, and ω_{12} are measured values, while E_J, E_C and g_{01} are extracted by solving the generalized Jaynes-Cummings Hamiltonian constrained by measured spectroscopic parameters.

which we can characterize with the framework of cQED. At low input microwave power (~ -120 dBm, Fig. 4.5(a) red/light gray traces) the cavity resonant frequency is shifted by the presence of the transmon circuit in its ground state. In the dispersive limit of cQED [34], $|\Delta| = |\omega_c - \omega_{01}| \gg g_{01}$, where ω_{01} is the qubit ground-to-excited state frequency, this hybridization causes the cavity resonant frequency to shift by an amount $\delta\omega \approx g_{01}^2/\Delta$. We measure $\delta\omega$ for both the empty cavity and the cavity full of superfluid helium, with the results reported in Table 4.1.

We utilize two-tone spectroscopy [100] to directly measure the excitation spectrum of the qubit and how it is modified by the superfluid. We use a low power tone (Fig. 4.5 (b), right of dashed line) to excite the qubit from ground $|0\rangle$ to first excited state $|1\rangle$, and a higher power tone to excite a two photon transition from $|0\rangle$ to the second excited state $|2\rangle$ (Fig. 4.5 (b), left of dashed line). From these measurements, we extract the $|0\rangle \rightarrow |1\rangle$ transition frequency ω_{01} and the $|1\rangle \rightarrow |2\rangle$ transition frequency ω_{12} for both the empty and full cavity configurations, and report these values in Table 4.1.

To extract the values of E_J, E_C , and g_{01} , and how they are modified by the dielectric

superfluid, we diagonalize the generalized JCH, and fit the eigenvalues ω_{01} , ω_{12} , and $\delta\omega$ to the values obtained from our spectroscopy measurements for the case when the 3D cavity is empty as well as when it is filled with helium. The results are summarized in Table 4.1.

The small change of E_J in the presence of helium is consistent with variations in E_J that we observe between cool downs without helium present in the cavity. It has been reported that these variations result from changes in the microscopic charge configuration in the Josephson junction oxide barrier [151, 134]. Therefore our results are consistent with E_J being unmodified by the presence of liquid helium. In contrast we find that the capacitive charging energy of the qubit decreases by 0.82%. This reduction in E_C agrees with the value of 0.78% obtained from finite element simulations of the system.

A shift in the vacuum Rabi coupling g_{01} is also induced by the superfluid helium. Qualitatively, this shift results from a change in the zero point energy of the cavity *and* a spatial redistribution of electric field lines within the cavity/qubit system upon changing the dielectric constant from $\epsilon = 1 \rightarrow \epsilon_{He} = 1.057$. Quantitatively, we may write the vacuum Rabi coupling in terms of the fluctuating zero point voltage of the microwave field in the 3D cavity [34] V_{ZPF} ,

$$g_{01} = 2eV_{ZPF}\beta\langle 1|\hat{n}|0\rangle, \quad (4.9)$$

where $\hat{n}$ is the Cooper pair number operator, and β is a parameter describing the efficiency of converting voltage fluctuations in the cavity to voltage fluctuations across the junction of the qubit [34]. We develop a simple model to capture the change of the vacuum Rabi coupling g_{01} as a function of the dielectric constant of the environment surrounding the qubit. Modeling the cavity as a simple LC oscillator (see Fig. 4.6) allows us to write the

zero point fluctuations of the voltage in the cavity as [152]

$$V_{ZPF} = \omega_c \sqrt{\frac{\hbar Z_c}{2}} \quad (4.10)$$

where $\omega_c = 1/\sqrt{L_c C_c}$ is the resonant frequency and $Z_c = \sqrt{L_c/C_c}$ is the impedance of the oscillator. Uniformly filling the cavity with a dielectric will shift the cavity capacitance from C_c to ϵC_c , and from Eqn. 4.10 one finds that

$$V_{ZPF} \propto \epsilon^{-3/4} \quad (4.11)$$

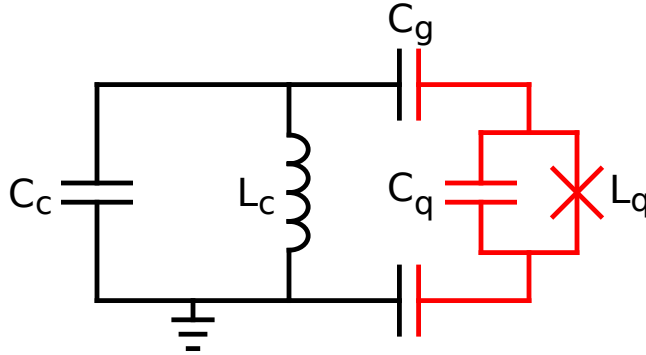


Figure 4.6: **Simplified circuit model for finding Δg_{01} :** Circuit model for a 3D transmon circuit (red/gray) coupled to a linear cavity (black) used to estimate the change in the vacuum Rabi splitting Δg_{01} in the presence of superfluid helium.

To understand the functional dependences of β for our experiment, we model the qubit as a parallel capacitance⁸ C_q and nonlinear Josephson inductor L_q coupled to the cavity via capacitance C_g (see Fig. 2). We assume the system is symmetric and that both of the antenna paddles of the qubit are identical and have the same capacitance C_g to the 3D cavity walls. β is then given by the voltage that builds up across C_q when a voltage V exists across the entire circuit,

⁸This capacitance includes both the intrinsic Josephson junction capacitance and the shunt capacitance provided by the antenna paddles of the qubit.

$$\beta = \frac{C_g^2}{C_g^2 + 2C_q C_g} \quad (4.12)$$

In the case of uniform dielectric filling, the capacitances will all scale uniformly and β will be unchanged from its vacuum value. We note, however, that the presence of the silicon chip and the intrinsic Josephson junction capacitance will cause C_g and C_q to scale differently as a function of dielectric constant of the cavity medium. We perform finite element simulations using COMSOL to determine how these capacitances change in the presence of helium. We find that $\Delta C_q = 0.78\%$ upon filling the cavity with helium, which agrees very well with the measured shift of 0.82% extracted from the change in the charging energy E_C , and that $\Delta C_g = 1.65\%$.

We finally note that the transmon excitation number transition matrix element is proportional to the zero point charge fluctuations of the qubit, which in the nearly harmonic oscillator regime of the transmon circuit may be written as

$$|\langle j+1 | \hat{n} | j \rangle| \approx \sqrt{\frac{j+1}{2}} \left(\frac{E_J}{8E_C} \right)^{1/4} \propto E_C^{-1/4} \propto C_q^{1/4} \quad (4.13)$$

We use the simulated shift in C_q to calculate the expected change of the charge number matrix element. Combining the shift in the matrix element with the predicted shifts in V_{ZPF} and β , we arrive at a predicted shift in the vacuum Rabi coupling induced by the superfluid in the cavity

$$\Delta g_{01} = -3.2\%, \quad (4.14)$$

which is in good agreement with our measured value of -2.8% .

4.4.2 Qubit Relaxation and Decoherence

Successful integration of qubits into quantum fluid experiments or vice versa requires that the coherence properties of the qubit do not degrade when immersed in superfluid helium. To characterize these effects, we use standard pulsed techniques to measure the energy (T_1) and phase (T_2) relaxation of the qubit as a function of temperature. We use standard measurement techniques, described in detail in Chapter 3 and in Fig. 4.7, to measure T_1 and T_2 . When measuring T_2 , the pulses are detuned from ω_{01} by ~ 300 kHz, allowing us to extract both the phase relaxation time and the qubit resonant frequency ω_{01} (see § 3.4.4 for details.) Additionally we conducted preliminary echo experiments for both the empty and superfluid-filled cavity configurations at the lowest temperature and found that the spin echo time T_2^e was not significantly different than T_2 .

To account for long timescale fluctuations of the qubit decoherence, we repeat the measurements of both T_1 and T_2 over a span of 5 hours. In Fig. 4.7(c-d) we plot the energy relaxation time T_1 , the pure dephasing time $T_\phi = (1/T_2 - 1/2T_1)^{-1}$ and $|\omega_{\text{drive}} - \omega_{01}|$ as a function of time for the empty (c) and superfluid filled (d) cavity configurations at the lowest temperature of the dilution refrigerator (10 mK). From Fig. 4.7(c-d) it is clear that the coherence times fluctuate significantly over the span of several hours (see Fig 4.7(e-f)). To account for these fluctuations while extracting decoherence information comparable between the empty and full cavity configurations, at every temperature examined we take data for 5 hours, and take the average and standard deviations of 5 hours of measurement to make a single datapoint in Fig 4.8. We plot the temperature dependences of T_1 (Fig. 4.8(a)), T_2 (Fig. 4.8(c), circles) and T_ϕ (Fig. 4.8(c), triangles) for the case when the cavity is empty (open symbols) and when it is filled with helium (closed symbols). The data points in Fig. 4.8

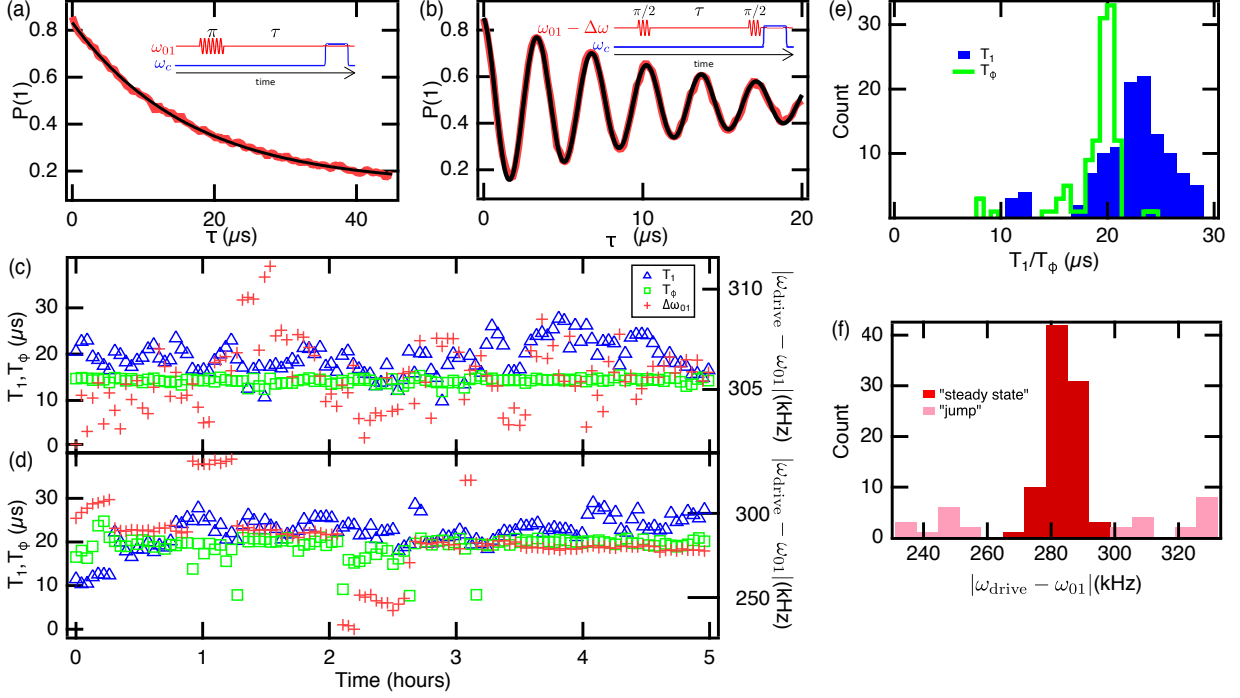


Figure 4.7: **Description of time-resolved and averaged measurements:** (a) Representative measurement of the qubit energy relaxation time T_1 inset with a schematic of the corresponding pulse sequence. We measure the probability $P(1)$ of finding the qubit in the excited state $|1\rangle$ after a variable delay time τ after exciting it and fit the data to an exponential function to extract the decay time T_1 . (b) A representative free induction decay measurement, which is fit to a sinusoid superimposed on a decaying exponential function. From this fit we extract the dephasing time T_2 and the drive/qubit detuning $|\omega_{\text{drive}} - \omega_{01}|$. We interleave a single run of each measurement described in (a) and (b), and then repeat this for ~ 150 s to get a single data set to fit to. We repeat this process for 5 hours to gather statistics on long timescale fluctuations: (c) is a representative measurement of $T_1, T_\phi = (1/T_2 - 1/2T_1)^{-1}$, and $|\omega_{\text{drive}} - \omega_{01}|$ for the empty cavity at $T \simeq 10$ mK. (d) Is a similar measurement run at the same temperature but for the cavity filled with helium. Note the difference in scale between right axes of (c) and (d). (e) A histogram of the values of T_1 and T_ϕ plotted in (d): The average and standard deviations of these datasets are what is reported in Fig (4.8). (f) A histogram of the values of $\Delta\omega_{01}$ recorded in (d): to measure the frequency shift as a function of temperature, we reject data measured during a discrete “jump” (light) and average over only data in “steady state” values of ω_{01} (dark).

represent the average value of a set of repeated measurements, while the error bars are the standard deviation of each set.

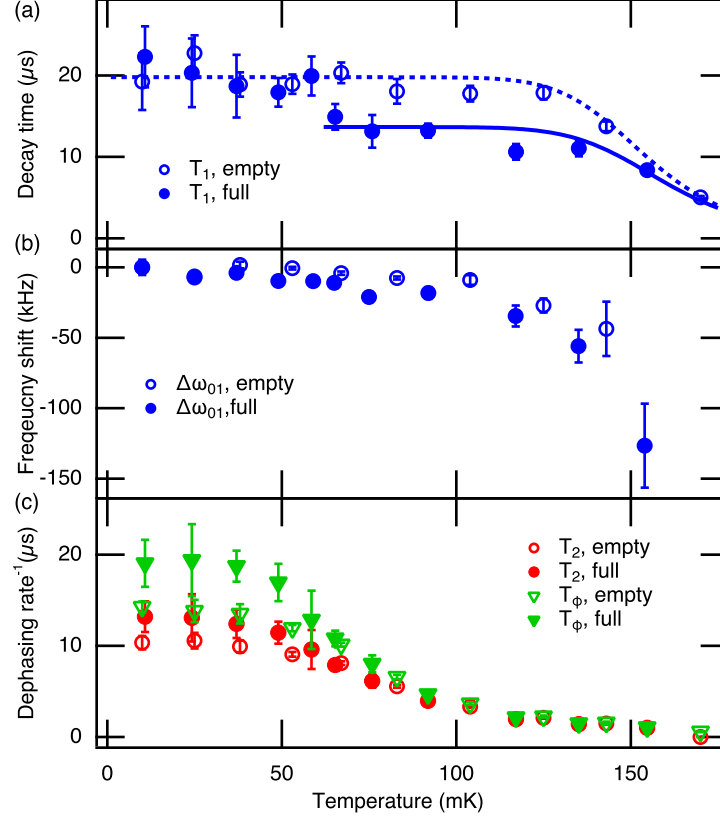


Figure 4.8: T_1 , T_2 and ω_{01} as a function of temperature with/without superfluid helium: (a) Qubit energy relaxation time T_1 , as a function of temperature, for both the empty cavity and the cavity filled with superfluid helium. We fit the data to theory (see Ref. [153]) for quasiparticle limited T_1 , using only the superconducting gap of aluminum and the nonequilibrium quasiparticle density x_{qp} (solid and dashed curves) to extract the change in quasiparticle density above ~ 60 mK when the cavity is filled with superfluid (see main text for discussion.) (b) Qubit frequency shift relative to its base temperature value, (c) dephasing (T_2) and pure dephasing (T_ϕ) times of the qubit as a function of temperature, for both the empty and full cavity configurations.

4.4.2.1 Qubit energy relaxation

At the lowest temperatures we find that T_1 saturates at roughly the same value ($\sim 20 \mu\text{s}$) both when the cavity is empty and when it is full of superfluid helium. Whatever mechanism is limiting T_1 at the lowest temperature we can conclude that it is not significantly suppressed

by the presence of superfluid helium.⁹ On the other hand, this result also demonstrates that the superfluid does not introduce any significant additional mechanisms for qubit energy relaxation at the temperatures relevant to cQED.

The temperature dependence of T_1 when the cavity is not filled with helium (see Fig. 4.8(a)) may be understood as arising from quasiparticles tunneling across the qubit junction [139, 153, 154, 142], limited by an athermal quasiparticle bath below ~ 140 mK. However, when the cavity is filled with helium we observe a qualitatively different temperature dependence: as we increase the temperature above ~ 60 mK, we find a modest reduction in T_1 when the cavity is filled with superfluid. We posit that this reduction in T_1 could be associated with a higher nonequilibrium quasiparticle density when the cavity is filled with helium. It is known that quasiparticles may travel long distances between superconducting islands on a substrate via conversion into phonons [145]. It is possible that at intermediate temperatures, phonons in the superfluid helium may mediate the transfer of quasiparticles between the superconducting qubit and the superconducting *cavity*. If the cavity were to have a higher nonequilibrium quasiparticle density, this additional coupling channel via the superfluid could cause the quasiparticle density in the qubit, and the associated relaxation rate via quasiparticle poisoning, to increase. At lower temperatures this additional source of quasiparticles would diminish as the Kapitza boundary resistance between the helium and the cavity/qubit continues to increase [89, 155]. This would lead to an increase in T_1 with decreasing temperature that would ultimately be limited by the same source that is limiting T_1 in the case of the empty cavity.

⁹Note that at the large qubit/cavity detuning used in this experiment, the increase in the Purcell emission rate Γ_p caused by the helium induced shift of the cavity frequency is negligible compared to the long timescale fluctuations in the emission rate. Specifically $\Gamma_p = (g_{01}/\Delta)^2 \kappa \sim (265 \mu\text{s})^{-1}$ for the empty cavity and $\sim (240 \mu\text{s})^{-1}$ for the full cavity, where κ is the cavity linewidth [34].

Working within this hypothetical model, we partially fit, down to ~ 60 mK, the temperature dependent T_1 data for the case when the cavity is full of helium to the theoretical quasiparticle decay rate given in [153] (see solid curve in Fig. 4.8(a)). In this fit the only parameters are the normalized nonequilibrium quasiparticle density x_{qp} and the superconducting gap of aluminum Δ_{BCS} . We find $\Delta_{BCS} \simeq 160 \mu\text{eV}$, which agrees well with the known bulk value for aluminum, and that an increase in quasiparticle density of $\Delta x_{\text{qp}} = 4 \times 10^{-6}$ accounts for the observed difference in T_1 above 60 mK when the cavity is filled with superfluid.

If the helium is mediating the introduction of extra quasiparticles into the qubit at intermediate temperatures, this increased quasiparticle density should also cause a shift in the resonant frequency of the qubit [153] relative to its zero temperature value. An increase in quasiparticle density of $\Delta x_{\text{qp}} \sim 4 \times 10^{-6}$ would produce a shift of $\Delta\omega_{01}/(2\pi) \approx -14$ kHz in the qubit frequency. This quasiparticle induced shift in ω_{01} must, however, be deconvoluted from shifts in the qubit resonant frequency caused by other sources of noise. We observe discrete abrupt changes (“jumps”) in the qubit frequency ω_{01} of the order $25 - 50$ kHz that occur over $\sim$ hour timescales, similar to those reported in other cQED experiments. We commonly observe (see Fig. 4.7(c)) that ω_{01} will jump from some steady state value, stay at the new value for minutes to hours, and then return to the original steady state value. Discrete changes such as these are commonly attributed to critical current noise in the Josephson junction [40]. Importantly, we can rule out shifts in the quasiparticle density as the origin of these jumps in ω_{01} , as they would be accompanied by an associated shift in the qubit relaxation rate $\Gamma_1 = 1/T_1$ of order $25 - 12 \mu\text{s}^{-1}$ [153], which we do not observe.

Therefore, to extract the temperature dependent quasiparticle induced shift in ω_{01} shown in Fig. 4.8(b), we employ a clustering algorithm to bin the measurements of ω_{01} around the

steady state frequency and reject measurements that occur during a discrete change in the qubit frequency. Fig. 4.7(f) shows an example of this for the data trace in Fig. 4.7(d): the central (dark) data are accepted while the outlying (light) data that was recorded during a discrete jump in ω_{01} is rejected. In Fig. 4.8(b) we plot the resonant frequency of the qubit as a function of temperature and find that the qubit frequency does in fact shift down appreciably in this intermediate temperature regime when the cavity is filled with helium. While this data is consistent with our hypothesis of an increased quasiparticle density at intermediate temperatures, our current experiment cannot directly confirm this model of superfluid phonon mediated quasiparticle coupling between the qubit and the 3D cavity. We will discuss how one may further confirm or refute this hypothesis in § 4.5.

Finally, while the temperature dependence of our T_1 data can be predominantly understood from the perspective of athermal quasiparticle poisoning, there likely are other mechanisms affecting the qubit energy relaxation. In particular the presence of near resonant two-level system (TLS) defects [136, 137, 138] can be a potential source of decreased T_1 . Further discussion on these potential mechanisms may be found in §4.4.2.3.

4.4.2.2 Qubit dephasing

In contrast to the energy relaxation of the qubit, we find that above 60 mK the pure dephasing time T_ϕ is the same both when the 3D cavity is empty and when it is full of superfluid. Upon cooling below ~ 60 mK we find that the dephasing time modestly improves in the presence of helium, indicating that qubit dephasing and energy relaxation are dominated by different mechanisms. Experiments similar to ours are known to be plagued by thermal photon occupations well above the nominal temperature of the mixing chamber of the dilution refrigerator [118, 88, 87, 90], and fluctuations in the cavity photon number have

been identified as a major limitation to the phase coherence of transmon qubits [76, 88, 90]. Additionally, in these previous experiments (see particularly [76, 87, 88]) the temperature dependence of the dephasing rate is qualitatively similar to that which we observe in our measurements both with and without helium.

It is known that many components in the microwave circuit are inefficiently thermalized, and it has been suggested that dissipative components such as attenuators may heat the microwave environment within the cavity to temperatures well above the dilution refrigerator temperature [87, 90]. It is possible that the superfluid helium is serving to better thermalize the microwave environment within the cavity in our experiment. For example, the helium could be opening an additional channel to cool the microwave circuitry via the central pins of microwave coupling lines or by directly cooling the 3D cavity walls. The modest improvement we observe in qubit dephasing when our cavity is filled with superfluid would correspond to a relatively minor reduction in the thermal photon number in the cavity. In the limit $\kappa \ll \chi$ where κ is the cavity linewidth and χ is the shift in the qubit frequency per cavity photon, and the limit $\bar{n}_{th} \ll 1$ where $\bar{n}_{th}$ is the thermal cavity photon occupancy, we can express the dephasing rate arising from residual thermal photons in the cavity as [156, 90]

$$\Gamma_{\phi} = \bar{n}_{th} \kappa \chi^2 / (\kappa^2 + \chi^2) \quad (4.15)$$

Using this expression we can estimate the temperature of the photon bath $T_{ph} \sim 80$ mK for the empty cavity and $T_{ph} \sim 70$ mK for the superfluid filled cavity.¹⁰ Finally, we note that these results are reproducible over multiple cool-down cycles of the cryostat.

¹⁰At both these temperatures, $\bar{n}_{th} \ll 1$, justifying the use of Eqn. 4.15

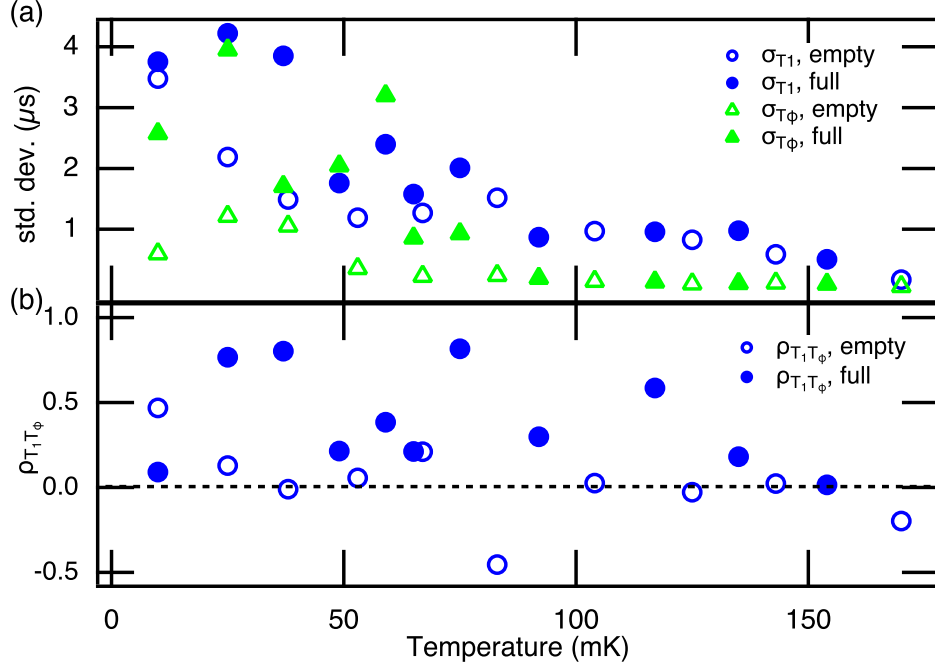


Figure 4.9: **Long timescale fluctuations of coherence parameters with/without superfluid helium:** (a) Standard deviation of the measured datasets of T_1 and T_ϕ for both empty and full cavities. (b) Normalized covariance of the sets of measured T_1 and T_ϕ .

4.4.2.3 Long timescale fluctuations in qubit coherence properties

In addition to qubit dephasing produced by a fluctuating cavity photon number, several dephasing mechanisms are known to be important to superconducting qubits. Single junction transmon qubits, like ours, are known to be insensitive to both charge and flux noise [76], and the dephasing caused by quasiparticle poisoning is predicted to be negligible compared to the relaxation induced by quasiparticles (i.e. $\Gamma_{\phi, \text{qp}} \ll \Gamma_{1, \text{qp}}/2$ [157].) Another possible source of dephasing is from TLSs nearly resonant with the qubit: as these TLSs undergo spectral diffusion, the associated dispersive shifts will also fluctuate, leading to qubit dephasing. Fluctuations of near resonant TLSs can also create a fluctuating density of states into which the qubit can decay. If TLSs are a dominant source of both transverse and longitudinal noise we should expect some correlation between the fluctuations in T_1 and T_ϕ .

In Fig. 4.9(b) We plot the normalized covariance of T_1 and T_ϕ

$$\rho_{T_1, T_\phi} = (\langle T_1 T_\phi \rangle - \langle T_1 \rangle \langle T_\phi \rangle) / (\sigma_{T_1} \sigma_{T_\phi}) \quad (4.16)$$

where $\rho_{T_1, T_\phi} = 1$ corresponds to perfectly correlated values of T_1 and T_ϕ and $\rho_{T_1, T_\phi} = 0$ corresponds to T_1 and T_ϕ being completely uncorrelated. We find that when the cavity is empty there is little correlation between T_1 and T_ϕ , and that this is broadly true also for the case when the cavity is filled with helium. There are, however, several measurement sets when the cavity is full of superfluid helium where $\rho_{T_1, T_\phi} \approx 0.8$, indicating strong correlation between T_1 and T_ϕ during these measurements. However, the more pronounced systematic increase in T_ϕ when the qubit is immersed in helium taken along with the fact that the correlation between T_1 and T_ϕ does not show any clear systematic temperature dependence indicates that, while perhaps not a dominant mechanism, TLSs could be playing a relatively larger role in the qubit energy decay and dephasing in the presence of superfluid helium for at least a subset of the measurements.

In addition, we also observe significant long timescale fluctuations of the decay and decoherence times T_1 and T_ϕ , which are not associated with changes we observe in ω_{01} . Several recent studies [136, 137, 138, 135] have attributed the long timescale fluctuations in T_1 to TLSs in proximity to the qubit both spectrally and in real space. As these TLSs fluctuate in frequency, they potentially provide a time-varying density of states into which the qubit can lose energy. In Fig. 4.9 we plot the standard deviation σ_{T_1} (σ_{T_ϕ}) of each set of T_1 (T_ϕ) measurements as a function of temperature and see that, generally, at lower temperatures the magnitude of the fluctuations is increased when the cavity is full of helium. These results indicate that unsaturated TLS fluctuators could be playing a role in qubit decoherence at the

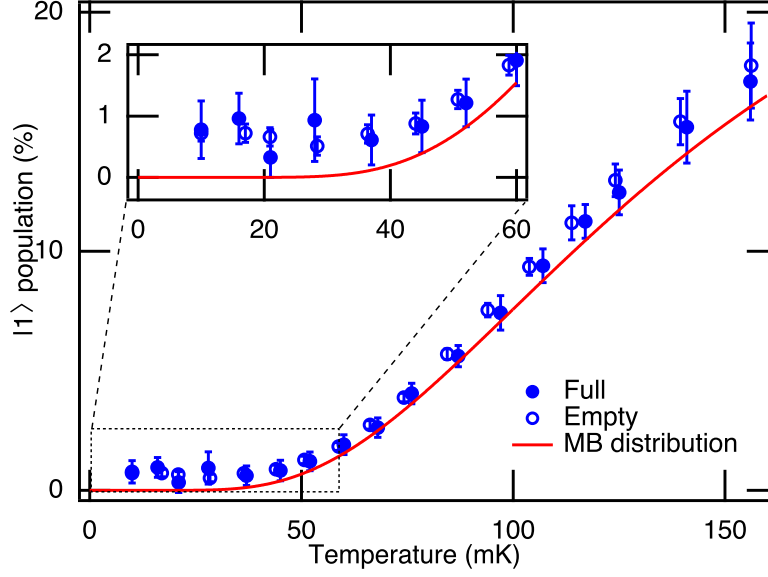


Figure 4.10: **Excited state population of the qubit as a function of temperature with/without superfluid helium:** Residual population of the qubit excited state $|1\rangle$ as a function of temperature for both the empty and superfluid filled cavity configurations. Also plotted is the theoretical Maxwell-Boltzman (MB) distribution, calculated using the energy levels obtained from spectroscopy. The presence of helium in the cavity has no significant effect on the $|1\rangle$ population, and the data fit the expected population well with no adjustable parameters. Inset: expanded view of boxed region. The $|1\rangle$ population for the qubit saturates at roughly the same value both when the cavity is full of superfluid and when it is empty.

lowest temperatures of our experiment when the cavity is full of helium. However, further experiments optimized for spectral [137, 138] or time domain [135] analysis of these fluctuations are needed to confirm the role that liquid helium has on TLS thermalization/fluctuation.

4.4.3 Residual excited state population

To further investigate possible thermalizing effects produced by the superfluid helium we have directly measured the residual qubit excited state population using a method developed in Refs. [158, 154] (see Appendix B for more details.) The measured population is plotted as a function of temperature in Fig. 4.10, along with the expected population calcu-

lated from a Maxwell-Boltzmann distribution with a partition function truncated beyond the 3rd excited state of the qubit. As shown in Fig. 4.10, the data are in good agreement with the theoretical population calculated with no adjustable parameters. Apparently, the superfluid helium has no significant effect on the residual excited state population of the qubit, which saturates at 0.5% – 1% at the lowest temperature in both the empty and full cavity configurations. This is consistent with the known difficulty of effectively cooling ^4He in the low milli-Kelvin temperature range due to the Kapitza thermal boundary resistance [89, 110] between superfluid helium and solid materials. Additionally, recent experiments attribute the majority of the residual excited state population to athermal quasiparticle poisoning [142]. The saturation of the qubit excited state at roughly the same value independent of the presence of superfluid helium is therefore also consistent with our measurements of T_1 (which we find to also saturate at roughly the same value) being mainly limited by athermal quasiparticles at low temperatures.

4.5 Discussion and future work

The three separate decoherence metrics (dephasing, depolarization, and residual excited state population) measured in this experiment seem to point to the same conclusion: in our experiment, the superfluid helium is only able to assist thermalization of the system down to $\sim 60 - 70$ mK. The depolarization time T_1 is more or less independent of temperature below $T_{MXC} = 60$ mK, indicating the mechanism causing extra depolarization in the presence of superfluid is frozen out at below that temperature. The pure dephasing time increases at the lowest temperatures, and attributing the dominant mechanism of dephasing to shot noise in the microwave cavity yields an effective cavity temperature of $T_{ph} \sim 70$ mK in the

presence of superfluid versus an effective temperature $T_{ph} \sim 80$ mK in the absence. This indicates that the superfluid is able to thermalize the cavity mode slightly more efficiently, however the cavity mode is decoupling from T_{MXC} at ~ 70 mK. The transmon residual excited state population saturates at a temperature below this $60 - 70$ mK threshold even without the superfluid: thus, the lack of significant change is consistent with the idea that the thermalization mediated by the superfluid largely decouples from the system degrees of freedom at these temperatures.

When thinking about future experiments one could pursue in studying superconducting qubits in the presence of superfluid helium, the data presented here opens up several interesting questions that could potentially be addressed with minor adjustments to the experimental apparatus/protocols. The helium seems to thermally decouple from the system at $60 - 70$ mK, much higher than the base temperature of the dilution refrigerator: is it possible to modify the experiment such that a lower passive cooling temperature is reached? Additionally, the experiment performed here is unable to say much about the *mechanisms* of decoherence, and how the superfluid impacts these mechanisms. Could we modify the experiment to gain information on the different mechanisms of decoherence, and how passive thermalization mediated by the superfluid impacts them? There are several future experimental steps one could imagine doing to address these questions:

Better superfluid thermalization: This experiment was not necessarily set up for optimal thermalization. The superconducting aluminum resonator acts as a thermal open circuit, and the nearest sinter heat exchanger (on the MXC plate) is separated from the cavity by nearly a meter of stainless steel capillary line. An obvious step one could take is to use a *copper* microwave cavity, which will more easily thermalize to the mixing chamber plate and help to locally thermalize the helium. Another step one could take is to put a sinter

heat exchanger in the experimental cell, providing a large surface area for the superfluid to thermalize to locally. Since the sinter heat exchanger would functionally act as a powder filter, to avoid microwave loss the sinter should be placed in a volume close to the 3D cavity but where the cavity mode has minimal amplitude. An example design implementing these changes is shown in Fig. 4.11(a).

However, the utility of ^4He for passive thermalization may be finite: since there are no free electrons in insulating ^4He , all thermal conductivity between helium and a solid must come from phonons in the liquid converting to phonons in the solid. The slow speed of sound, as well as the phonon population that decreases as T^{-3} , yields an increasing Kapitza thermal boundary resistance of ^4He as temperature decreases [89], limiting its utility as a thermalizing agent at low temperatures. On the other hand, ^3He is known to have a superior thermal boundary resistance at very low temperatures, likely because the nuclear magnetic moment of ^3He allows for the exchange of magnetic excitations between the solid and the fluid [89]. A superconducting qubit device passively cooled by liquid ^3He would likely thermalize to a much lower temperature, however such an experiment would require a substantially different design: The aforementioned magnetic moments of ^3He could drive decoherence in the system, especially if a flux-tunable qubit were employed, and the large volume of the 3D cavity in this experiment makes simply reproducing the experiment using ^3He prohibitively expensive.

Further understanding decoherence mechanisms: In recent years, interest has increased in understanding the long timescale fluctuations of the coherence properties of superconducting circuits, in both the time [136, 137, 138] and frequency [135] domains. Long timescale fluctuations have been attributed to two-level systems [136, 137, 138] and quasiparticles bursts from stray background radiation [144, 143], however a general consensus

of the limiting factors of decoherence is not necessarily clear. From this perspective, it would be useful to have an *in situ* knob, i.e. the presence or absence of thermalization via the superfluid, that may influence these mechanisms in different ways.

The measurements in the papers cited above largely rely on measuring the decoherence properties of a superconducting qubit system with a time resolution of ≤ 1 s, either to raster through qubit frequency in a reasonable time or to build up a power spectral density of coherence property fluctuations over several orders of magnitude in frequency. In the data presented in this chapter, during a coherence time measurement run, each *individual* measurement in the five-hour average, corresponding to a single point in Fig.4.7(c), takes about 120 s. This time is a function of the experimental rep-rate and how many sequences we average to get an individual point, and could be decreased significantly with a higher rep-rate¹¹ and less averaging. In our lab, we have demonstrated the ability to get reasonable quality data to fit to and extract decoherence times (as well as precise measurements of ω_{01}) in as little as 1-2 seconds. Therefore, it should be fairly straightforward to systematically investigate the effect of superfluid helium by measuring the spectrum of decoherence fluctuations over several orders of magnitude, or by using a flux tunable qubit in a copper cavity to raster the qubit in frequency space.

Testing the quasiparticle limited T_1 hypothesis: In addition to gleaning information about decoherence mechanisms from fluctuations in the coherence properties of a qubit, it is possible to lightly modify the setup and directly measure the tunneling of superconducting quasiparticles across the junction. This has been observed using an “offset charge sensitive” (OCS) transmon, with $E_J/E_C \sim 15 - 25$ reduced from the normal transmon op-

¹¹Note that a high experimental rep-rate may artificially increase the residual $|1\rangle$ population, as the qubit may still be excited from the previous sequence. A good rule of thumb is $(\text{rep rate})^{-1} > 5 \max(T_1)$.

erating regime [140, 142, 42]. In this regime, ω_{01} retains a small but measurable dependence on the offset charge n_g : since a single superconducting quasiparticle carries charge $q = 1e$, a tunneling event will abruptly shift n_g by $1/2$, shifting the qubit frequency in turn (see Fig. 2.3.) The quasiparticle parity (i.e. whether a tunneling event has occurred or not) may then be measured by measuring the evolution of the qubit state-vector referenced to a tone at the average value of the ω_{01} band (see Ref. [142] for more details.)

From these measurements, one can infer parity jumps, i.e. quasiparticles tunneling events, in real time and infer the density of quasiparticles in the metal. Using this technique, one could unambiguously measure whether or not the superfluid has any effect on the quasiparticle density. One major drawback of this technique is that it requires high-fidelity single-shot readout of the qubit state, likely more sensitive than our current measurement setup. This could be accomplished by adding a quantum limited amplifier, such as a Josephson parametric amplifier, to the amplification chain, however doing so would require modest upgrades to our setup/experimental capabilities.

4.5.1 Coupling to superfluid acoustic modes?

To conclude the chapter, we note that the *increase* in pure dephasing time at the lowest temperatures is indicative that, at least in the setup tested here, there was very little in the way of coupling to mechanical modes in the superfluid. This result, in retrospect, isn't so surprising. In the experiment presented, we made no attempt to optimize the geometry of the system to preferentially couple to a specific mechanical mode of the system. One may imagine an experiment where this geometry is optimized: for example, if the transmon capacitance is provided by a vacuum gap capacitor, as sketched in Fig. 4.11(b), the electric field mode profile would have a strong overlap with the fundamental acoustic mode of helium trapped between

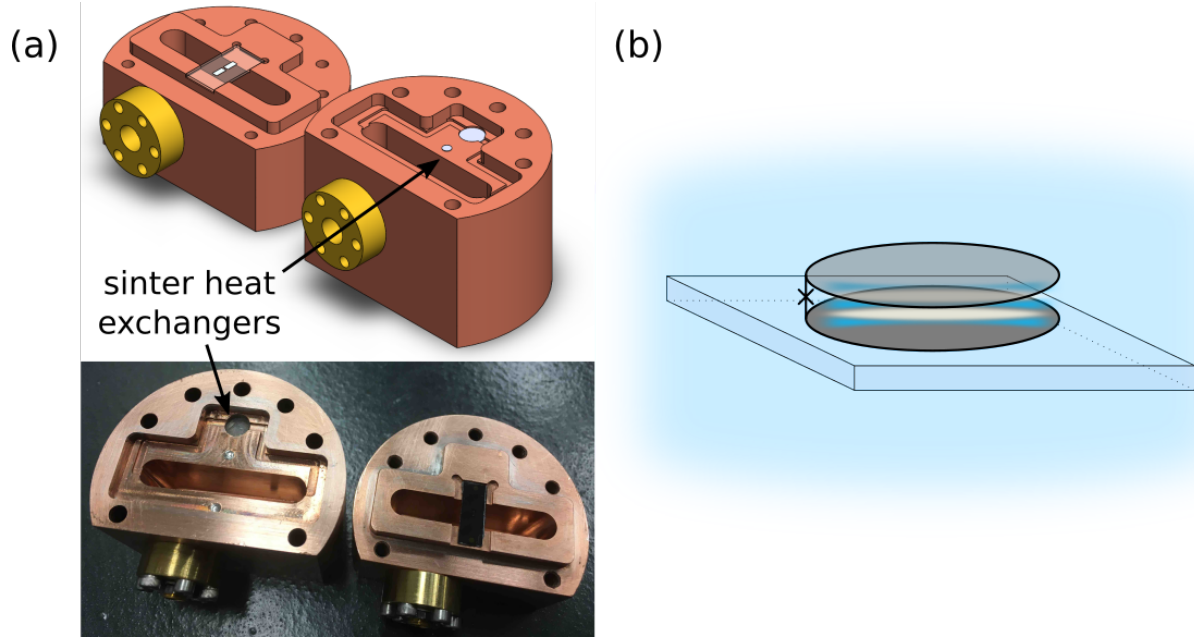


Figure 4.11: **Proposals for future superfluid-superconducting qubit experiments:** (a) Top: CAD model of an upgraded experimental cell for improving thermalization/studying the effects of superfluid helium on qubit decoherence. A copper cavity with local sinter heat exchangers would both better thermally anchor the helium to the cryostat temperature and allow for flux tunability of the qubit. Bottom: picture of the upgraded experimental cell, built for future experiments. (b) Sketch of a possible experiment optimizing the overlap of a qubit mode with an acoustic mode in the superfluid.

the capacitor plates. The envisioned geometry, with appropriate spacing/size to both localize a superfluid mode in the 2 – 5 GHz range and provide the correct transmon capacitance, should be achievable with standard vacuum gap capacitor fabrication techniques.

However, even this optimized geometry may not provide enough coupling to the acoustic mode in the helium. In the absence of piezoelectricity, fluctuations in the helium density cannot directly induce voltage fluctuations across the capacitor pads, as the zero point fluctuations of the electromagnetic cavity we considered in Chapter 2 do. The zero point fluctuations of the helium density will, however, modify the effective dielectric constant of the system, coupling the mechanical mode to the electrical mode. This situation is, in some sense, similar to the optomechanical coupling between a mechanical mode and an optical

resonator, where shifts in the optical cavity resonant frequency are caused by displacements of the mechanical oscillator. While not entirely analogous, cavity optomechanics experiments coupling the motion of helium to an electromagnetic cavity have reported single photon coupling rates¹² ranging from $g_0/(2\pi) = 10^{-8}$ Hz for a kHz acoustic resonator overlapping with an GHz microwave resonator [159, 110] to $g_0/(2\pi) = 3 \times 10^3$ Hz for a MHz acoustic resonator overlapping with an optical resonator [113]. To truly enter the regime of quantum acoustics, the vacuum Rabi coupling rate g_m between the qubit and the mechanical mode should satisfy $g_m > \max[\kappa_m, \Gamma_2]$, where κ_m is the loss rate of the acoustic cavity. While a low coupling rate does not preclude experiments demonstrating the coupling of a qubit to a highly populated mechanical mode, the low single photon coupling rates reported in superfluid optomechanics experiments make the prospect of superfluid quantum acoustics mediated *only* by electrostriction dim. To perform quantum acoustics experiments, we likely need a method to more strongly couple mechanical strain to electric fields, such as piezoelectricity.

¹²Optomechanics experiments may afford to have much lower single photon coupling, since many photons may be circulating in the linear optical resonator. The strong nonlinearity of qubits heavily limits the number of photons in the “optical” (qubit) mode, usually to one.

Chapter 5

Surface acoustic waves and surface acoustic wave devices

As indicated in our discussion of superfluid helium, experiments investigating single-excitation quantum acoustics with superfluid helium seem extremely challenging. In order to couple to acoustic modes in a substrate, we need a stronger method of mediating the electromechanical coupling, such as piezoelectricity. How should we go about coupling a quantum circuit to a mechanical mode using the piezoelectric effect?

Surface acoustic waves (SAWs) are elastic waves that propagate along the surface of an elastic material, with a displacement amplitude that decays exponentially away from the surface. In a piezoelectric material, as we shall see, SAWs may easily be excited by applying an ac voltage at frequency f_0 to a periodic structure with periodicity $\lambda = v/f_0$ where v is the propagation speed of SAWs. Since $v \sim 10^3 - 10^4$ m/s for typical SAW materials, at telecommunication frequencies (1-10 GHz), SAW devices have a characteristic length scale $\sim 0.1 - 10 \mu\text{m}$. These dimensions are easily achievable with modern lithographic techniques, and moreover the fact that the structures are confined to the *surface* make them straightforward to fabricate: often times only a single layer of metal needs to be deposited on the surface.

The short wavelength ($10^4 - 10^5 \times$ shorter than the wavelength of light at the same

frequency), and ease of fabrication have made SAW devices ubiquitous in modern telecommunication electronics. Common components made from SAWs include delay lines, bandpass filters, and oscillators, and have found applications spanning from gas sensors to garage door openers [160, 161].

5.1 Elastic waves in solids

Before we discuss piezoelectric surface waves, it will behoove us to recall some properties of waves in elastic materials. This discussion largely follows that of [160], with some help from [162]. We'll start off by thinking about an isotropic elastic medium: a “particle”¹ that sits at rest in its equilibrium position $\vec{r}$ may undergo some displacement $\vec{u}(\vec{r}, t)$ from $\vec{r}$ (see Fig. 5.1.) When we talk about elasticity, we are interested in *deformations* of the material, not overall displacements or rotations (which cause no internal forces and can be simply mapped onto coordinate changes.) We may satisfy these conditions by constructing the strain tensor at a point $\vec{r}$

$$S_{ij}(\vec{r}, t) = \frac{1}{2}(\partial_j u_i + \partial_i u_j) \quad (5.1)$$

where $\partial_i = \partial/\partial r_i$. From this form, we can read off that S_{ij} is invariant under overall displacements (constant offsets in $\vec{u}$.) We also see that S_{ij} is a symmetric tensor, implying that there are only 6 independent components of the strain tensor.²

Having assembled an object that describes displacement caused by internal forces, we

¹Here, we're considering our elastic medium in the continuum limit: the “particle” is much smaller than the length scale of any relevant deformations (i.e., the wavelength of the acoustic waves we'll consider) but much larger than a single unit cell of the underlying crystal.

²In the literature, this symmetry is often taken advantage of to write S_{ij} as a 6 component vector $S_{ij} = (S_{11}, S_{22}, S_{33}, S_{23}, S_{13}, S_{12})^T$. This representation is called the Voigt notation.

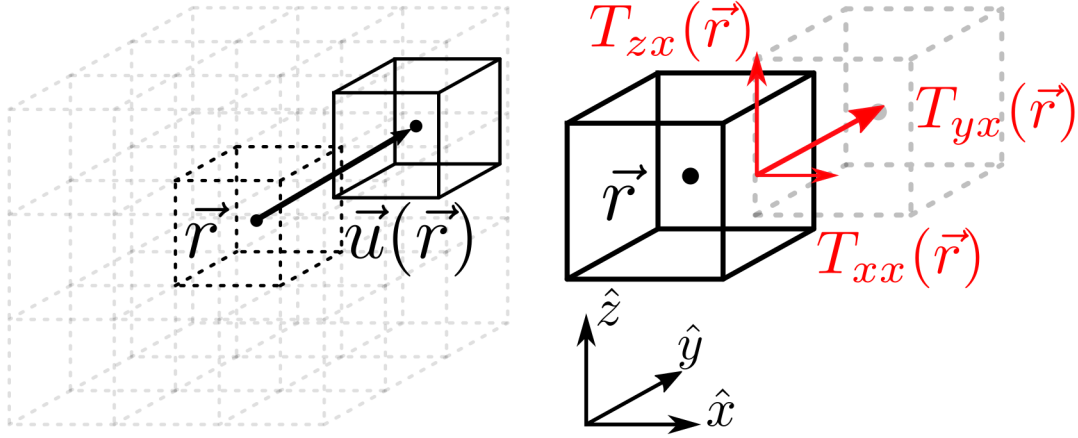


Figure 5.1: **Stress and strain in elastic media:** Left: to understand stress and strain in an elastic medium, we conceptually break up the medium into “particles” located at point $\vec{r}$ when the system is in equilibrium (dashed black cube.) When forces are applied to the system, the particle at point $\vec{r}$ is displaced from its equilibrium by $\vec{u}(\vec{r})$ (solid black cube, greatly exaggerated.) Right: A defining feature of elastic media is that adjacent “particles” may exert forces on each other when displaced from equilibrium. The neighbor (dashed) adjacent to the $\hat{x}$ surface of the particle at $\vec{r}$ (solid) may be displaced in any direction, which will exert a force on the particle at $\vec{r}$ in the displacement direction. The force per unit area this particle exerts in the $\hat{i}$ direction is T_{ix} .

now must describe those internal forces: the stress tensor $T_{ij}(\vec{r}, t)$ accomplishes this goal. Consider our “particle” of elastic material as a cube: the side of this cube that is normal to the $\hat{x}$ direction experiences a force caused by displacements in the adjacent material (see Fig. 5.1.) Crucially, this force may be in any direction: the T_{ix} component of the stress tensor is defined as the force per unit area in the $\hat{i}$ direction acting on the surface of our particle normal to $\hat{x}$. Thus, diagonal components of the stress tensor represent compressive forces, while off-diagonal components represent shear forces. It may also be shown that $T_{ij} = T_{ji}$, and that subsequently the stress tensor also only has 6 independent components. We may relate stress and strain using a generalized version of Hooke’s law

$$T_{ij} = c_{ijkl} S_{kl} \quad (5.2)$$

where c_{ijkl} is the rank-4 stiffness tensor, and summation over repeated indices is implied.³

To write down the equations of motion for the material, it's convenient to again think of our “particle” as a cube, this time with explicit edge length δ centered at $\vec{r} = (x_0, y_0, z_0)$. Consider the two faces of the cube at $x = x_0 \pm \delta/2$. Since T_{ix} is the force per unit area in the $\hat{i}$ direction on surfaces normal to $\hat{x}$, the total force on the cube from these surfaces is

$$\delta^2 \left(T_{ix}(x_0 + \delta/2, y_0, z_0) - T_{ix}(x_0 - \delta/2, y_0, z_0) \right) = \delta^3 \frac{\partial T_{ix}(x_0, y_0, z_0)}{\partial x}$$

where, in the second equality, we have pulled out a factor of $1/\delta$ and gone to the limit $\delta \rightarrow 0$. The forces acting on the surfaces normal to the $\hat{y}$ and $\hat{z}$ directions will take similar forms, and will simply add with the forces on $\hat{x}$ surfaces. Thus, the total force on the cube in the $\hat{i}$ direction is

$$F_i = \delta^3 \partial_j T_{ij}$$

Since the mass of the particle is $\rho \delta^3$, where ρ is the material's mass density, we can write down the Newtonian equations of motion as

$$\rho \ddot{u}_i = \partial_j T_{ij} = \frac{1}{2} \partial_j c_{ijkl} (\partial_k u_l + \partial_l u_k) \quad (5.3)$$

where $\ddot{u} = \partial^2 u / \partial t^2$.

Up until now, our discussion has been completely general: these equations are valid for small displacements in any elastic material. It will be instructive for us to break from this

³Similar to S_{ij} , T_{ij} is also commonly written down in the Voigt notation as a 6-component column vector. We can infer from the symmetries of S_{ij} and T_{ij} that c_{ijkl} will, at most, have 36 independent components, since in Voigt notation c_{ijkl} can be written a 6×6 matrix. In fact, it can be shown that this 6×6 matrix itself is symmetric, meaning c_{ijkl} has at most 21 independent components. The underlying crystal symmetry of the material will almost always constrain these components further.

generality and consider the propagation of waves in *isotropic* materials, i.e. materials whose properties have no dependence on the orientation of the coordinate frame we choose. This is a curious choice, since by definition piezoelectric materials are *anisotropic* (they break inversion symmetry), however there is much intuition to gain from briefly studying waves in isotropic materials.

The stiffness tensor of an isotropic material may be written down as [160]

$$c_{ijkl} = \nu \delta_{ij} \delta_{kl} + \mu (\delta_{ik} \delta_{jl} + \delta_{il} \delta_{jk}) \quad (5.4)$$

where ν and μ are the Lamé coefficients of the material⁴ and δ_{ij} is the Kronecker delta. Plugging this form of the stiffness tensor into Eqn. 5.2 yields, after some simple tensor algebra

$$T_{ij} = \nu \delta_{ij} S_{kk} + 2\mu S_{ij} \quad (5.5)$$

We can plug this form into the Eqn. 5.3 to derive the wave equation for waves in an isotropic material

$$\begin{aligned} \rho \ddot{u}_i &= \partial_j (\nu \delta_{ij} S_{kk} + 2\mu S_{ij}) \\ &= \partial_j (\nu \delta_{ij} \partial_k u_k + \mu (\partial_i u_j + \partial_j u_i)) \\ &= \nu \partial_i \partial_k u_k + \mu \partial_i \partial_j u_j + \mu \partial_j^2 u_i \\ \rho \ddot{\vec{u}} &= (\nu + \mu) \vec{\nabla} (\vec{\nabla} \cdot \vec{u}) + \mu \vec{\nabla}^2 \vec{u} \end{aligned} \quad (5.6)$$

While this is not exactly the familiar wave equation, it does relate the second time derivative

⁴It is conventional in the literature to use λ and μ as the Lamé coefficients, however to avoid confusion with the SAW wavelength I have switched notation to ν and μ .

of u to the second spatial derivative of u . It is therefore a reasonable guess that, in the absence of boundary conditions (i.e. in an infinite isotropic medium), solutions should take the form of plane waves

$$\vec{u} = \vec{u}_0 e^{-i(\omega t - \vec{k} \cdot \vec{r})} \quad (5.7)$$

Plugging this form of the solution into Eqn. 5.6 yields an algebraic equation for the amplitudes of waves in isotropic materials

$$\omega^2 \rho \vec{u}_0 = (\nu + \mu)(\vec{k} \cdot \vec{u}_0) \vec{k} + \mu |\vec{k}|^2 \vec{u}_0 \quad (5.8)$$

From this, we can pick out two types of elastic waves that may propagate in solids. The first case is if $\vec{k} \cdot \vec{u}_0 = 0$, i.e. the displacement field is *transverse* to the wave vector. In this case, the dispersion relationship of the wave is given by

$$|\vec{k}_t|^2 = \omega^2 \rho / \mu \quad (5.9)$$

and we may read off the velocity of these transverse (or shear) waves as $v_t = \sqrt{\mu/\rho}$. There are two possible polarizations of transverse waves, in the two directions perpendicular to the propagation direction. For example, if we take the wave to be propagating in the $\hat{z}$ direction, we may write down a general form for transverse waves as $\vec{u}_t(r, t) = (u_x \hat{x} + u_y \hat{y}) e^{-i(\omega t - kz)}$. Since $\vec{\nabla} \cdot \vec{u}_t = 0$, it is easy to check that for transverse waves, Eqn. 5.6 reduces to the standard wave equation with speed of sound v_t :

$$\ddot{\vec{u}}_t = v_t^2 \vec{\nabla}^2 \vec{u}_t \quad (5.10)$$

A second type of wave is possible when $\vec{k} \cdot \vec{u}_0 \neq 0$. In order to have a nontrivial solution to Eqn. 5.8 satisfying this condition, the wave vector and the displacement field must be parallel to each other. In that case, $(\vec{k} \cdot \vec{u}_0)\vec{k} = |\vec{k}|^2 \vec{u}_0$ and the dispersion relationship of these waves is given by

$$|\vec{k}_l|^2 = \omega^2 \rho / (2\mu + \nu) \quad (5.11)$$

and we may read off the velocity of these longitudinal (or compression) waves as $v_l = \sqrt{(2\mu + \nu)/\rho}$. Since μ and ν are always positive [160], longitudinal waves in an isotropic material always propagate faster than transverse waves. Once again, if we take the wave to be propagating in the $\hat{z}$ direction, we may write down a general form for longitudinal waves as $\vec{u}_l(r, t) = u_z \hat{z} e^{-i(\omega t - kz)}$. Noting that the $\vec{\nabla} \times \vec{u}_l = 0$ and remembering the vector identity $\vec{\nabla}(\vec{\nabla} \cdot A) = \vec{\nabla}^2 A + \vec{\nabla} \times \vec{\nabla} \times A$, we once again see that Eqn. 5.6 reduces to the familiar form of the wave equation, albeit it with a different speed of sound v_l :

$$\ddot{u}_l = v_l^2 \vec{\nabla}^2 \vec{u}_l \quad (5.12)$$

5.1.1 Rayleigh Waves

In the previous section, we discussed the two types of plane waves (longitudinal and transverse) that may propagate in the bulk of an isotropic medium, where we have not defined any boundary conditions. The surface of the substrate will impose a boundary condition, which will constrain the types of waves that may propagate: namely, since there is no material above the surface to exert forces on the material at the surface, the stress on the surface normal to $\hat{z}$ must be zero, i.e. $T_{xz} = T_{yz} = T_{zz} = 0$. As we will see shortly, SAWs are in fact

a linear superposition of transverse and longitudinal waves. These surface waves are named after Lord Rayleigh who first described them 1885 [163], and we will use the terms “Rayleigh wave” and “SAW” interchangeably. This derivation largely follows that of [164], and [160].

We’ll start off by building up a reasonable guess for the general form of a SAW, and then use the boundary conditions to solve for a specific solution. To get started, imagine the situation presented in Fig. 5.2(a), where an elastic material fills space $z < 0$ with vacuum in $z > 0$. Along the surface at $z = 0$, a monochromatic surface plane-wave propagates in the $\hat{x}$ direction. We will assume for this discussion that we have complete translational symmetry along the $\hat{y}$ direction⁵, and that the displacement field of the SAW obeys both $u_y = 0$ and $\partial_y \vec{u} = 0$.

Our strategy will be to decompose the SAW displacement field into the sum of its transverse and longitudinal components, which each individually obey wave equations 5.10 and 5.12 respectively, i.e. $\vec{u} = \vec{u}_l + \vec{u}_t$. Since we are seeking a monochromatic plane-wave solution, both the longitudinal and transverse components should propagate the same way in the $\hat{x}$ direction, and we may write down a general form for the functional dependence of each individual component of the wave as

$$|\vec{u}_{l,t}| \propto e^{-i(\omega t - \beta x)} f_{l,t}(z) \quad (5.13)$$

where β and ω are the same for both the longitudinal and transverse components, and the z -dependence $f_{l,t}(z)$ depends on the polarization. Inserting each component into its respective

⁵SAWs with displacement perpendicular to both the propagation direction and the vector normal to the surface do exist, and are called shear-horizontal waves. However, they will not concern us in this dissertation.

wave equation (Eqn. 5.10 or 5.12) yields a differential equation for the z -dependence:

$$\partial_z^2 f_{l,t}(z) = \left(\beta^2 - \frac{\omega^2}{v_{l,t}^2} \right) f_{l,t}(z)$$

from which we infer the functional form of $f_{l,t}(z)$:

$$f_{l,t}(z) = e^{\pm \kappa_{l,t} z}, \quad \kappa_{l,t} = \sqrt{\beta^2 - \frac{\omega^2}{v_{l,t}^2}} \quad (5.14)$$

If the wave is to be confined to the surface, we must impose the requirement that $\beta^2 - \omega^2/v_{l,t}^2 > 0$, or else the amplitude will oscillate into the bulk. Note that, since we expect the velocity of the SAW to obey $v = \omega/\beta$, this condition implies $v < v_{l,t}$, i.e. that SAWs will propagate slower than either bulk transverse or bulk longitudinal waves. Since, in our setup, the elastic material occupies half space $z < 0$, we should take the $+\kappa_{l,t}$ solution, else the amplitude of the displacement will blow up at $z \rightarrow -\infty$.

Having ironed out the functional dependence of $\vec{u}_{l,t}$, we may now write down a specific vector solution for each component. We start with the transverse component: recall that, for transverse elastic waves, $\vec{\nabla} \cdot \vec{u}_t = \partial_x u_{t,x} + \partial_z u_{t,z} = 0$. Inserting the solutions 5.13 and 5.14 yields a constraint on the relative amplitudes of the $u_{t,x}$ and $u_{t,z}$

$$i\beta u_{t,x} + \kappa_t u_{t,z} = 0$$

from which we may write down the transverse component of the SAW wave up to an overall constant A

$$\vec{u}_t = A(\kappa_t \hat{x} - i\beta \hat{z}) e^{-i(\omega t - \beta x) + \kappa_t z} \quad (5.15)$$

Likewise, the longitudinal component of the SAW will satisfy the condition $\vec{\nabla} \times \vec{u}_l = \partial_z u_{l,x} - \partial_x u_{l,z} = 0$. Using this condition, we insert solutions 5.13 and 5.14 to constrain the relative amplitudes of the $u_{l,x}$ and $u_{l,z}$

$$\kappa_l u_{l,x} - i\beta u_{l,z} = 0$$

from which we may write down the longitudinal component of the SAW wave up to an overall constant B

$$\vec{u}_l = B(\beta\hat{x} - i\kappa_l\hat{z})e^{-i(\omega t - \beta x) + \kappa_l z} \quad (5.16)$$

We are now in a position where we can invoke the boundary conditions provided by the surface, i.e. that the stress on the free surface must be zero. We may use Eqn. 5.5 to write down three differential equations, one for each boundary condition:

$$T_{xz}|_{z=0} = \partial_z u_x + \partial_x u_z = 0 \quad (5.17a)$$

$$T_{yz}|_{z=0} = \partial_z u_y + \partial_y u_z = 0 \quad (5.17b)$$

$$T_{zz}|_{z=0} = v_l^2 \partial_z u_z + (v_l^2 - 2v_t^2) \partial_x u_x = 0 \quad (5.17c)$$

where in Eqn. 5.17c I have dropped the y -dependence and written down μ and ν in terms of the longitudinal and transverse velocities v_l and v_c . Clearly Eqn. 5.17b is satisfied by our decision to work with displacement fields obeying $u_y = 0$ and $\partial_y \vec{u} = 0$. We may then insert the total displacement field $\vec{u} = \vec{u}_l + \vec{u}_t$ (where $\vec{u}_l$ and $\vec{u}_t$ are given by 5.16 and 5.15 respectively) into 5.17a and 5.17c to generate a system of linear equations for the amplitudes A and B . After a bit of algebra, we arrive at the following set of homogeneous

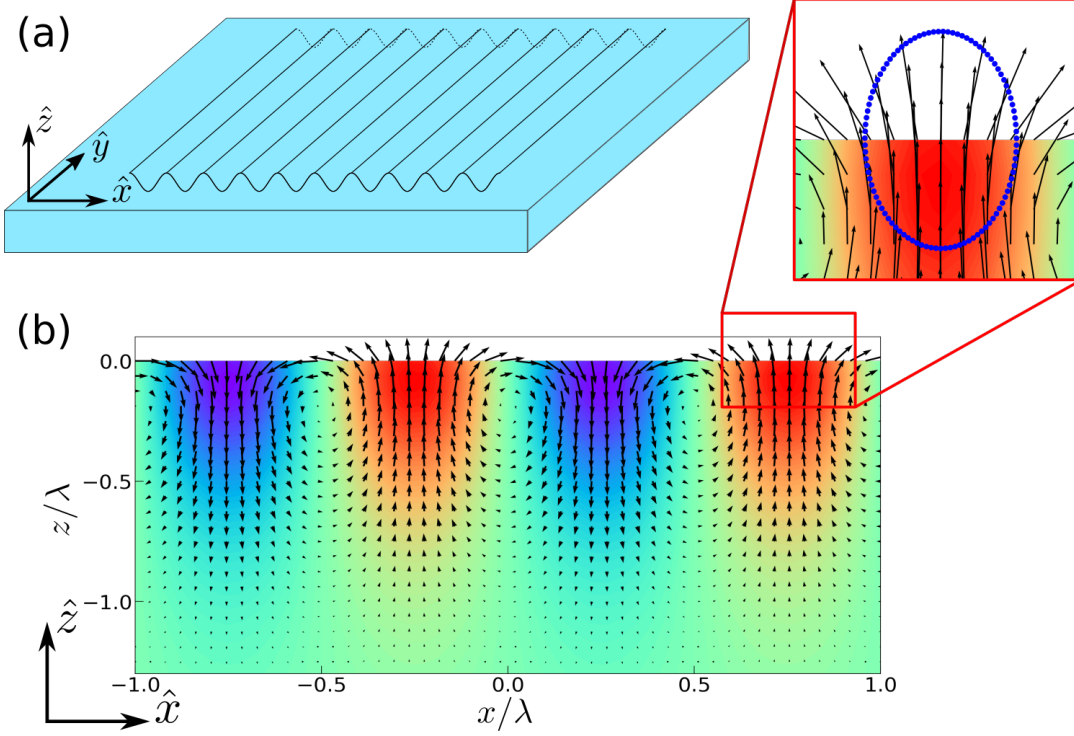


Figure 5.2: **SAW propagation in an isotropic medium:** (a) Schematic of SAW propagation: we search for surface wave solutions where the surface is normal to the $\hat{z}$ -direction and the SAW propagates in plane waves along the $\hat{x}$ -direction. (b) Plot of the SAW displacement field $\vec{u}(r)$ for $v_t = v_l/\sqrt{3}$. The length/direction of the arrows corresponds to the displacement amplitude/direction respectively, and color corresponds to real part of the vertical displacement u_z . Inset: The (amplitude exaggerated) trajectory of a single “particle” as a function of time. Over one period of SAW propagation, the particle traces out an ellipse. This figure was inspired in part by Fig. 2.1 of Ref. [165].

linear equations:

$$\begin{aligned} A(\kappa_t^2 + \beta^2) + B(2\beta\kappa_l) &= 0 \\ A(2\beta\kappa_t) + B(\kappa_t^2 + \beta^2) &= 0 \end{aligned} \tag{5.18}$$

In order for a nontrivial solution to this set of equations to exist, the determinant of the matrix composed of the coefficients must be equal to zero. Writing the SAW velocity as $v = \omega/\beta$ and using the condition 5.14, this gives us an equation for the SAW velocity

$$\left(2 - \frac{v^2}{v_t^2}\right)^2 = 4\sqrt{1 - \frac{v^2}{v_t^2}}\sqrt{1 - \frac{v^2}{v_l^2}}$$

which becomes more illuminating when we write down the SAW velocity as traverse velocity in the bulk times some constant of proportionality ξ , i.e. $v = \xi v_t$

$$(2 - \xi^2)^2 = 4\sqrt{1 - \xi^2}\sqrt{1 - \xi^2\left(\frac{v_t}{v_l}\right)^2} \quad (5.19)$$

From this, we see that the SAW velocity in an isotropic material is independent of frequency: it is in fact only dependent upon the ratio v_t/v_l . While squaring this equation gives a polynomial equation with multiple roots, we demand that ξ be both positive and real, and that $\xi < 1$ (otherwise $v > v_t$, and waves would propagate into the bulk.) These conditions yield one solution for ξ , which turns out to be $\xi \approx 0.875 - 0.955$ for possible values of v_t/v_l [164]. Thus, SAWs propagate with a velocity slightly lower than that of transverse waves in the bulk. We may also use Eqn. 5.18 to solve for the ratio of the amplitudes $A/B = (\xi^2 - 2)/(2\sqrt{1 - \xi^2})$. Since $\xi < 1$, this ratio is real, and the longitudinal and transverse components of the SAW oscillate in phase with each other. Inspecting Eqns. 5.15 and 5.16, we see that this implies that the $\hat{x}$ and $\hat{z}$ components of the SAW displacement oscillate *out of phase* with each other, and a “particle” on the surface will follow an elliptical path. The full solution for $\vec{u}(r)$, with proper relative amplitudes (up to an overall constant) is plotted in Fig. 5.2(b). The blue dots in the inset track the trajectory of a single “particle”, indeed tracing out an ellipse as the SAW undergoes a single oscillation.

5.2 Piezoelectricity and piezoelectric Rayleigh waves

In certain materials, the application of strain results in the generation of a macroscopic electric field, and conversely an external electric field may induce strain in the material. This effect, known as the piezoelectric effect, happens in anisotropic materials whose underlying structure breaks inversion symmetry, and was first discovered by Jacques and Pierre Curie in 1880 [160]. In a non-piezoelectric medium, the electric displacement $\vec{D}$ may be related to the electric field $\vec{E}$ using the familiar formula $D_i = \epsilon_{ij}E_j$, where ϵ_{ij} is the dielectric tensor of the medium. In piezoelectric materials, we must tack on an additional term to this relationship, as strain in these materials causes macroscopic polarization fields. Assuming the strain/field are small enough such that we need to only consider linear terms, the electrical displacement field in a piezoelectric material may be written down as

$$D_i = \epsilon_{ij}E_j + e_{ijk}S_{jk} \quad (5.20)$$

Here, e_{ijk} is the rank-3 piezoelectric tensor, which relates the rank-2 strain S_{ij} to the rank-1 displacement field⁶. The converse to this is that electric field may induce stress in the materials. This is accounted for by generalizing Eqn. 5.2

$$T_{ij} = c_{ijkl}S_{kl} - e_{kij}E_k \quad (5.21)$$

where e_{kij} is the same piezoelectric tensor as before⁷. It is crucial to remember that the piezoelectric effect occurs explicitly in anisotropic materials, and thus our previous conversation of waves in isotropic materials will not strictly apply here. In reality, the symmetry

⁶Keeping with the Voigt notation, it is common to write down e_{ijk} as a 3×6 matrix.

⁷In the Voigt notation, the index shuffling amounts to this term containing the transpose of the 3×6 piezoelectric matrix.

breaking of many crystals makes the analysis of propagating waves rather painful, and often times only solvable by numerical techniques. Finding an explicit solution for the propagation of piezoelectric SAWs is beyond the scope of this thesis, and would not contribute greatly to our understanding of the results anyway. That being said, we will be able to recycle much of the intuition built in the previous section to build a qualitative description of piezoelectric Rayleigh waves.

5.2.1 Equations of motion and Boundary conditions

Since sound in elastic materials propagates $\sim 10^5 \times$ slower than light, it is a good approximation to work with quasi-static electric fields, i.e. we may write the electric field as the gradient of the scalar electrostatic potential $\vec{E} = -\vec{\nabla}\vartheta$. This reduces the degrees of freedom we have to solve from 6 (3 components of $\vec{u}$ and 3 components of $\vec{E}$) to 4. This simplification allows us to write Newton's Law (Eqn. 5.3) for the motion of the deformation field in terms of the electrostatic potential

$$\begin{aligned}\rho\ddot{u}_i &= \partial_j T_{ij} = \partial_j (-e_{kij} E_k + c_{ijkl} S_{kl}) \\ &= e_{kij} \partial_j \partial_k \vartheta + \frac{c_{ijkl}}{2} (\partial_j \partial_k u_l + \partial_j \partial_l u_k) \\ \rho\ddot{u}_i &= e_{kij} \partial_j \partial_k \vartheta + c_{ijkl} \partial_j \partial_k u_l\end{aligned}\tag{5.22}$$

We will also assume that our substrate is an insulator, so that there are no free charges and $\vec{\nabla} \cdot \vec{D} = 0$. This provides another constraint to the system

$$\epsilon_{ij} \partial_i \partial_j \vartheta + e_{ijk} \partial_i \partial_j u_k = 0\tag{5.23}$$

Equations 5.22 and 5.23 define a system of 4 differential equations, which we may solve subject to the appropriate boundary conditions to find the fields $\vec{u}(r, t)$ and ϑ associated with elastic wave propagation in a piezoelectric material.

Similar to the isotropic case presented in §5.1.1, we will consider an elastic material that fills the half-volume $z < 0$, and search for a plane wave solution propagating in the $\hat{x}$ direction⁸ that is exponentially confined to the material surface $z = 0$. If the SAW has wave number β and oscillates at frequency ω , we can guess a general form of the functional dependence $\vec{u}$ and ϑ

$$\begin{aligned}\vec{u} &= \vec{u}_0 f(z) e^{-i(\omega t - \beta x)} \\ \vartheta &= \vartheta_0 f_\vartheta(z) e^{-i(\omega t - \beta x)}\end{aligned}\tag{5.24}$$

We will once impose the mechanical boundary conditions that the components of the stress $T_{xz} = T_{yz} = T_{zz} = 0|_{z=0}$. However, we have some flexibility with regards to the electrical boundary condition. There are two electrical boundary conditions that are important to consider:

1. A perfectly conducting surface. This situation may arise if, for example, a thin metal film is deposited onto the surface of the insulating piezoelectric substrate. In this case, the electrical boundary condition is $\vartheta|_{z=0} = 0$.
2. If the surface is also insulating, there is nothing stopping the electric field from extending into the vacuum above the surface. In vacuum, the potential obeys Laplace's equation $\vec{\nabla}^2 \vartheta = 0$, and we may insert the predicted form of the solution (Eqn. 5.24)

⁸Note that the coordinate system we define here is **not** associated the principle axes of the underlying crystal in any way. We choose a direction of SAW propagation, and rotate the coordinate system (and all the relevant material parameters c_{ijkl} , e_{ijk} and ϵ_{ij}) to align that direction with the $\hat{x}$ -axis.

into Laplace's equation to generate a differential equation for the z -dependence of the field in vacuum:

$$\partial_z^2 f_{\vartheta}(z) = \beta^2 f_{\vartheta}(z) \quad (z > 0)$$

This differential equation has solutions $f_{\vartheta}(z) = e^{\pm\beta z}$, and we choose the $-\beta$ solution so the electric field doesn't blow up at $z \rightarrow \infty$. Thus we arrive at a general feature of piezoelectric surface acoustic waves: in the absence of a conducting layer, the electrostatic potential (and thus electric field) associated with the wave **extends evanescently above the surface of propagation with a decay constant equal to the SAW wave number**. Equivalently, the decay length is equal to the SAW wavelength $\lambda/2\pi$.

Since there are no free charges on an insulating surface, we maintain that $\vec{\nabla} \cdot \vec{D} = 0$ everywhere, which implies that the $\hat{z}$ component of $\vec{D}$ must be continuous across the surface. Using the form of the solution above, this condition produces the boundary condition $D_z|_{z=0} = \beta\vartheta_0 e^{-i(\omega t - \beta x)}$.

5.2.2 Sketch of Piezoelectric Rayleigh waves

We are now in a position to sketch the solutions of Equations 5.22 and 5.23 that describe the propagation of piezoelectric SAWs. As stated previously, the anisotropic nature of piezoelectric materials makes general solution often only achievable through numerical methods, which are beyond the scope of this thesis. Therefore, we will only provide a qualitative overview of the solution, from which we will extract several useful concepts. It should be noted that there is an analytic solution for SAW propagation in materials in the cubic crystal system such as gallium arsenide, and a detailed derivation may be found in Ref. [166].

In §5.1.1, we constructed the displacement field of a Rayleigh wave as the sum of longitudinal and transverse components, each of which satisfy the equations of motion in the bulk. In the piezoelectric case, the same general prescription is followed: the deformation field and the electric potential are built out of *partial waves* with fixed frequency/wave number that satisfy the equations of motion in the bulk. To construct these partial waves, we start off with a reasonable guess for the functional form of the SAW fields

$$\begin{aligned}\vec{u} &= \vec{u}_0 e^{\gamma z} e^{-i(\omega t - \beta x)} \\ \vartheta &= \vartheta_0 e^{\gamma z} e^{-i(\omega t - \beta x)}\end{aligned}\tag{5.25}$$

Since all functional dependence of these fields are contained in exponents, the four differential equations 5.22 and 5.23 become a set of linear homogeneous equations for the amplitudes of the waves. These equations may be written as

$$W \begin{pmatrix} u_{x,0} \\ u_{y,0} \\ u_{z,0} \\ \vartheta_0 \end{pmatrix} = \begin{pmatrix} 0 \\ 0 \\ 0 \\ 0 \end{pmatrix}\tag{5.26}$$

where W is a matrix determined by the particular form of equations of motion (and, subsequently, by c_{ijkl} , e_{ijk} , and ϵ_{ij} .) In the same manner as with the isotropic case, we claim that for fixed wave number β , there should be some specific values of γ that allow for these equations to be satisfied. This is equivalent to saying that a nontrivial solution to Eqn. 5.26 requires $\det(W) = 0$. The determinant of W gives us an 8^{th} order polynomial in γ , with 8 roots in general, however we only take solutions where $\text{Re}(\gamma) > 0$, else the fields will blow up at $z \rightarrow -\infty$. Generally, 4 roots satisfy this condition, and thus the solution will be a

sum of 4 partial waves. Note that while we demanded $\text{Re}(\gamma) > 0$, unlike in the case of an isotropic material, γ may be complex, and thus the amplitudes of the displacement/electric potential may oscillate as well as decay into the bulk of the material.

For each root γ_m , there will be a particular solution to Eqn. 5.26 for the amplitudes $\vec{u}_0^{(m)}, \vartheta_0^{(m)}$, defined up to an overall constant. These parameters define the m^{th} partial wave, and a general solution is a superposition of the 4 partial waves

$$\begin{aligned}\vec{u} &= \sum_m A_m \vec{u}_0^{(m)} e^{\gamma_m z} e^{-i(\omega t - \beta x)} \\ \vartheta &= \sum_m A_m \vartheta_0^{(m)} e^{\gamma_m z} e^{-i(\omega t - \beta x)}\end{aligned}\tag{5.27}$$

where A_m is the relative amplitude of each of the 4 partial waves. In the same manner as in §5.1.1, we invoke the appropriate boundary conditions to generate another set of 4 homogeneous linear equations for the coefficients A_m , and obtain the velocity of the wave by setting the determinant of coefficients to zero.

5.2.3 Relevant parameters

Since piezoelectric materials are anisotropic, SAW propagation will depend on the coordinate system we choose (i.e., how are the surface, which we have taken to be normal to $\hat{z}$, and the SAW propagation direction, which we have taken to be $\hat{x}$, oriented with respect to the underlying crystallographic axes?) When designing SAW devices, our goal is often to maximize the coupling between electrical excitations and mechanical excitation. For a given material, how should we go about choosing how to orient the surface (also called the crystal cut, or simply cut) and the propagation direction?

As stated before, there are two electrical boundary conditions to consider: the case of an

Material	Cut	Propagation axis	v (m/s)	K^2 (%)
LiNbO ₃	Y	Z	3488	4.8
LiNbO ₃	128°-rotated Y	X	3979	5.4
Quartz	ST	X	3158	0.14
GaAs	(100)	[011]	2864	0.07

Table 5.1: Relevant material parameters for common SAW substrates, reproduced from Chapter 4 of Ref. [160] and Chapter 9 of Ref. [167], both of which contain more complete tables of materials/parameters.

insulating surface (where the electric field evanescently extends into the vacuum above the surface) and the case of a conducting surface (where the $\vartheta = 0$ at the surface.) When the procedure in §5.2.2 is executed subject to these two boundary conditions, it is typically found that a SAW propagating along an insulating surface propagates slightly faster than a SAW propagating along a conducting surface [160]. The relative size of this velocity shift is useful for quantifying piezoelectric coupling: we define the piezoelectric coupling coefficient for a given cut/propagation axis as $K^2 = 2\Delta v/v$, where v is the insulating surface velocity, v_m is the conducting (metalized) surface velocity, and $\Delta v = v - v_m$. For a given material, SAW cuts and propagation axes are typically chosen to maximize K^2 . Table 5.1 contains useful properties for some common SAW materials/cuts, reproduced from Refs. [167, 160]⁹. When working with SAW devices, it is conventional to specify both the cut and the propagation axis along with the material: for example, the quartz cut specified in Table 5.1 is commonly referred to as “ST-X Quartz”.

⁹Note that there seems to be some minor disagreement between these sources: when in conflict, we defer to Ref. [167]

5.3 SAW devices

We now turn to a description of electromechanical devices based on piezoelectric SAWs. As stated before, these devices are almost all periodic structures with periodicity λ fabricated on the surface of a piezoelectric substrate, and will operate at or around the so-called center frequency $f_0 = v/\lambda$. We have already established that a full solution of SAW propagation on a surface with translational symmetry is beyond the scope of this thesis, and breaking this symmetry/adding electromagnetic sources will only complicate the matter. Nevertheless, there exist powerful tools for describing the *electrical* response of a SAW device without directly solving the electromechanical equations of motion. Our strategy will be to use these techniques to describe the electrical response of a SAW device as a function of frequency, and map that response onto the response of simple circuit elements.

5.3.1 Interdigitated transducers

The most important component of any SAW device is the interdigitated transducer, or IDT. Interdigitated transducers are the method by which we convert electrical signals into acoustic signals (or vice versa): they consist of two electrodes, with interlocked metallic “fingers” of width a (see Fig. 5.3(a)) deposited on the surface of the substrate. The periodicity of the IDT is $\lambda = 2a + 2b$, where b is the spacing between the fingers. One wavelength will consist of a single finger from either electrode of the IDT, and so we generally characterize the length of the IDT by counting the number of finger pairs N_p . To keep with the convention from §5.1.1 where SAWs propagate in the $\hat{x}$ -direction, we will define the long axis of the IDTs as the $\hat{y}$ -direction: that way, when a voltage is applied across the two electrodes of the IDT the voltage is spatially periodic along the $\hat{x}$ -direction. The width of the finger overlap w defines

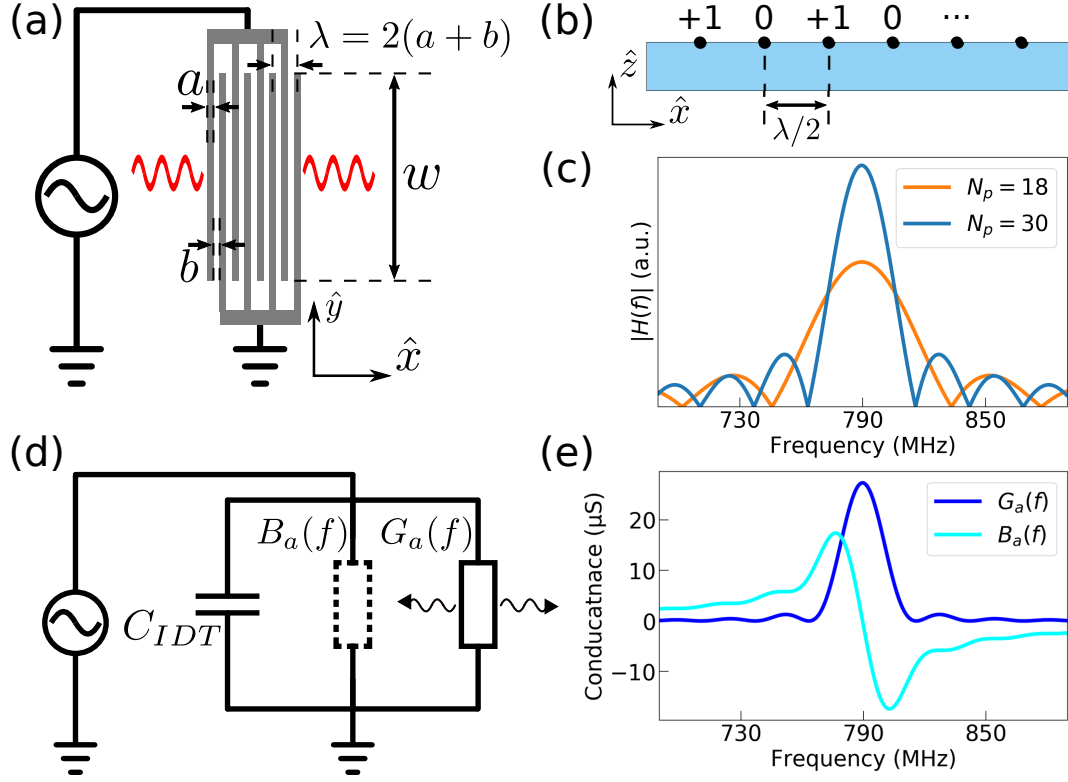


Figure 5.3: **Interdigitated transducers for SAWs:** (a) A schematic of an interdigitated transducer (IDT) of width w with periodicity $\lambda = 2(a + b)$. (b) Delta function model of infinitely thin IDT fingers, which yields the frequency response function like the one plotted in (c). (c) Frequency response calculated for two different IDTs on ST-X quartz, with $a = b = 1 \mu\text{m}$. As we increase the number of finger pairs of the IDT, the response at the center frequency f_0 ($\approx 790 \text{ MHz}$) becomes larger, and the bandwidth of the device becomes more narrow. (d) Full electrical model of an ideal IDT, including real ($G_a(f)$) and imaginary ($B_a(f)$) acoustic impedances, as well as the geometric capacitance C_{IDT} of the IDT structure. (e) Calculated $G_a(f)$ and $B_a(f)$ for the $N_p = 30$ device from (c), with $w = 60 \mu\text{m}$.

the width of the SAW beam excited by the IDT.

To understand the electrical response of an SAW IDT, let us approximate each finger of the IDT as an infinitely thin point source, evenly spaced along the $\hat{x}$ -direction [160, 168] (see Fig. 5.3(b).) We may then write down the applied voltage to the surface as a sum of delta functions, spaced by $\lambda/2$. If we hold one electrode at ground potential and apply a voltage $V = +1$ (in arbitrary units) to the other electrode, the applied voltage has an x -dependence

of

$$\vartheta \propto \sum_{n=0}^{N_p-1} \delta(x - \lambda n) \quad (5.28)$$

Fourier transforming this potential gets us the wave vector response (or transfer function)

$$H(k) \propto \sum_{n=0}^{N_p-1} e^{-i\lambda kn} = \frac{\sin(N_p k\lambda/2)}{\sin(k\lambda/2)} e^{-i(N_p-1)k\lambda/2} \quad (5.29)$$

Where to get the second equality we've used the identity $\sum_{n=0}^{N-1} r^n = (1 - r^N)/(1 - r)$ for a geometric series. Near the center frequency of the IDT, $f \approx f_0$, $k\lambda/2 \approx \pi$, and we may approximate $\sin(k\lambda/2) \approx -k\lambda/2$. Using the relationships $f_0 = v_s/\lambda$ and $k = 2\pi f/v_s$, we may write approximate the frequency response as a sinc function centered on f_0

$$H(f) \approx N_p \left| \frac{\sin(N_p \pi (f - f_0)/f_0)}{N_p \pi (f - f_0)/f_0} \right| \quad (5.30)$$

We can extract some valuable information from the form of this frequency response. Since the sinc function is maximized when the argument is 0, as expected, the most efficient conversion of electrical signals into SAWs (and vice versa) will happen at $f = f_0$. We also see that as we increase N_p , the magnitude of the response function will increase, while the bandwidth (the frequency range over which $|H(f)|/|H(f_0)|$ is of order unity) will decrease. Fig. 5.3(c) shows the calculated $|H(f)|$ plotted for two IDT's with different N_p .

We are now in a position to model the SAW IDT as a lumped element *electrical* component, and infer the frequency dependence of this component. When an oscillating voltage V is applied across the IDT, electrical power will be converted into propagating SAWs: we may define the real part of the IDT conductance via this dissipated power, i.e. $P = G_a |V|^2/2$.

Likewise, the power of a propagating wave is proportional to the square of the potential generated by the wave: $P_{SAW} \propto \vartheta_{SAW}^2$. In an ideal case, where all the power dissipated at the IDT is converted into SAWs, these two power are equal to each other, and we may infer that the frequency dependence of the real part of the acoustic conductance $G_a(f) \propto |\vartheta_{SAW}(f)|^2 \propto |H(f)|^2$. It may be shown that the acoustic admittance of the IDT also has an imaginary component $B_a(f)$ [167].

The constant of proportionality of G_a depends on the IDT geometry, and needs to be evaluated numerically [167, 160]. For $a = b$, this numerical evaluation leads to a full expression of the complex SAW IDT conductivity, modeled as 3 parallel lumped element components (see Fig. 5.3(d)):

$$Y_{IDT}(f) = G_a(f) + i(B_a(f) + 2\pi f C_{IDT}) \quad (5.31a)$$

$$G_a(f) = G_{a0} \left[\frac{\sin X}{X} \right]^2 \quad (5.31b)$$

$$B_a(f) = G_{a0} \left[\frac{\sin 2X - 2X}{2X^2} \right] \quad (5.31c)$$

where C_{IDT} is the unavoidable geometric capacitance of the IDT, and

$$\begin{aligned} X &= N_p \pi \frac{f - f_0}{f_0} \\ G_{a0} &\approx 1.3 K^2 N_p^2 (2\pi f) w C_s \end{aligned} \quad (5.32)$$

where C_s is the capacitance per IDT finger pair per unit length, and w, N_p, f_0 and K^2 are the aforementioned beam width, finger pair number, center frequency and piezoelectric

coupling coefficient respectively.¹⁰ An example calculation of $G_a(f)$ and $B_a(f)$ is plotted in Fig. 5.3(e).

5.3.2 Coupling of Modes and the P-matrix

The above discussion of ideal IDTs ignores the fact that the metal strips that form the IDT, in addition to transducing SAWs, may also *reflect* SAWs impinging upon them. The magnitude of the reflection coefficient $|r_s|$ of a single metal strip tends to be small ($\approx 0.1-2\%$, depending upon the substrate and design [160, 169, 15, 10]), and thus the above analysis is still fairly accurate if $N_p \ll |r_s|^{-1}$, where only a small fraction of the wave is reflected while in transit across the device. We will, however, encounter SAW devices that are explicitly designed to reflect SAWs, such as the Bragg mirrors that form a SAW cavity.

For devices where reflections are non-negligible, we use the coupling of modes (COM) method (see chapter 8 of Ref. [160] for a more complete discussion) to account for both transduction and reflections. The COM method, similar to the ideal transducer model above, assumes that transducers are delta functions that may both reflect and transduce SAWs. This model is used to set up a system of coupled differential equations relating the leftward propagating SAW amplitude A_L , rightward propagating SAW amplitude A_R , and oscillating current (I) and voltage (V) across the transducer. The system of equations is then solved subject to the appropriate boundary conditions to yield a 3-port scattering matrix, called the P-matrix [160, 167]

¹⁰Since $C_{IDT} = C_s N_p w$, we may easily compare the relative magnitude of the acoustic admittance and the capacitive admittance: at $f = f_0$, the ratio of the acoustic admittance to the capacitive admittance is $G_a(f_0)/2\pi f_0 C_{IDT} \approx 1.3K^2 N_p$. Thus we see that, the acoustic admittance doesn't start dominating until $N_p > (K^2)^{-1}$. For strongly piezoelectric substrates, such as LiNbO₃, this gives us a reasonable value $N_p \approx 20$, but for weakly piezoelectric substrates such as quartz or GaAs, the number of finger pairs required is absurdly large (500-1000), and fulfilling this requirement renders the IDT bandwidth unfeasibly small. Thus, on these substrates, the IDT response is typically dominated by the geometric capacitance.

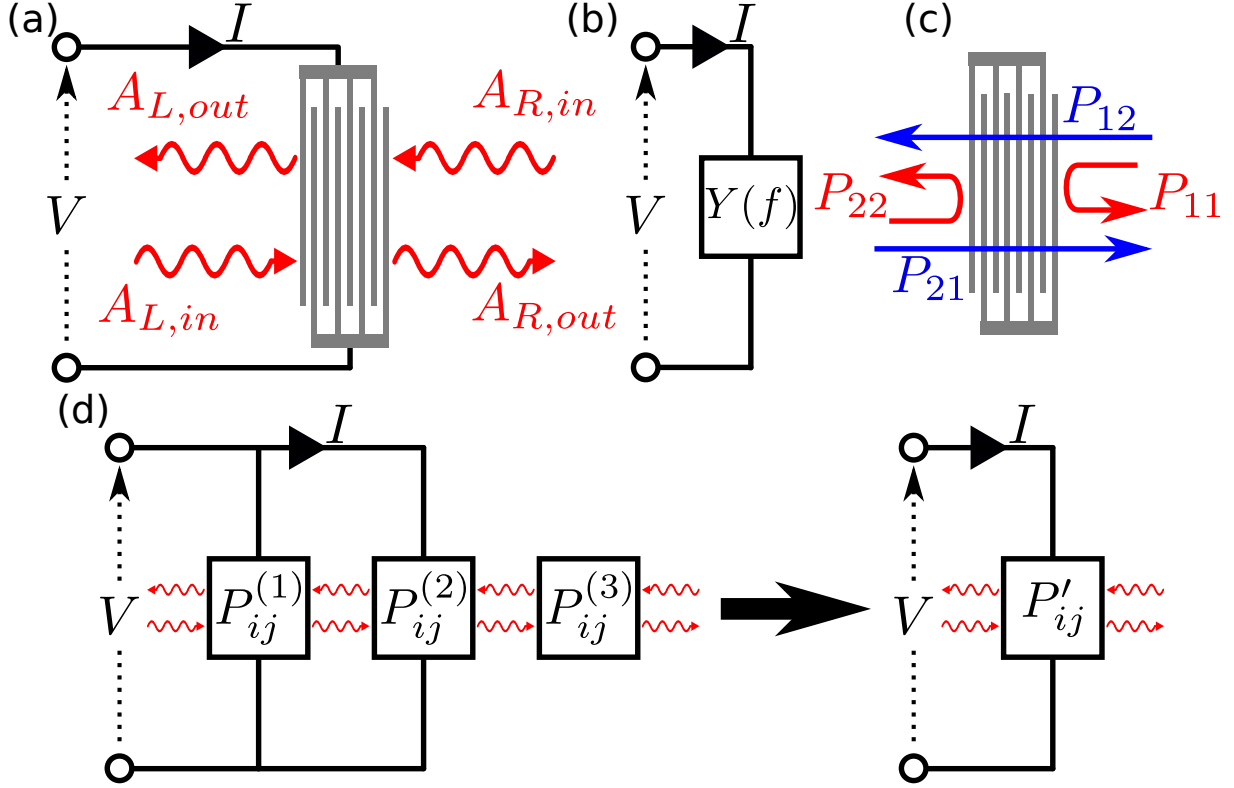


Figure 5.4: **The P-matrix for SAW devices:** (a) We may model a SAW element as a 3-port object, with two SAW ports (a left/right port with incoming SAW amplitude $A_{L,in}/A_{R,in}$ and outgoing amplitude $A_{L,out}/A_{R,out}$) and an electrical port with current I and voltage V . (b) In the absence of any incoming/outgoing SAWs, the only non-zero component of the P-matrix is P_{33} , which we recognize as the electrical admittance $Y(f)$. (c) In the absence of an electrical connection, the P-matrix collapses down to a two port scattering matrix describing the reflection of SAWs at the device. (d) The major advantage of the P-matrix formalism is that several devices may be cascaded into one composite P-matrix. These structures may be electrically connected in parallel (such as $P_{ij}^{(1)}$ and $P_{ij}^{(2)}$ in the example), however we may also cascade elements with no electrical connection (such as $P_{ij}^{(3)}$.)

$$\begin{pmatrix} A_{L,out} \\ A_{R,out} \\ I \end{pmatrix} = \begin{pmatrix} P_{11} & P_{12} & P_{13} \\ P_{21} & P_{22} & P_{23} \\ P_{31} & P_{32} & P_{33} \end{pmatrix} \begin{pmatrix} A_{L,in} \\ A_{R,in} \\ V \end{pmatrix} \quad (5.33)$$

Here, $A_{L,in}/A_{R,in}$ ($A_{L,out}/A_{R,out}$) are the amplitudes of the rightward/leftward SAW

propagating into (out of) the structure in question, as shown in Fig. 5.4(a). The P-matrix has several convenient features that are worth explicitly noting:

1. In the absence of SAWs ($A_{L,in}, A_{R,in} = 0$), $I = P_{33}V$. Therefore, $P_{33} = Y$ the electrical admittance of the transducer, which we will treat as a lumped element (Fig. 5.4(b).) In the end, P_{33} will be the electrical element we insert into our circuit model.
2. P_{11}, P_{12}, P_{21} and P_{22} describe scattering of SAWs off the structure. This sub-matrix can be thought of as a scattering matrix for SAWs along the path of a structure. For example, an ideal non-reflecting transducer will have $P_{12} = P_{21} = 1$ (up to a phase factor that the SAW picks up traversing the structure) and $P_{11} = P_{22} = 0$.
3. All the elements in the third row/column have to do with transduction (converting SAWs into an electrical signal/vice versa.) If we have a structure that doesn't transduce (such as a Bragg mirror that is electrically disconnected from the circuit), these elements are all zero, and the P-matrix will reduce to the SAW scattering sub-matrix.

The primary advantage of the P-matrix is that multiple structures may be *cascaded* with relative ease. This means that, in order to obtain P-matrix elements for a composite SAW structure such as a resonator (a launching IDT and two Bragg mirrors), rather than solving the COM equations for the entire structure, we only need the solution for each structure (Bragg mirrors and IDT) and can cascade the structures together to get a description of the entire device.¹¹ As such, we'll individually model Bragg mirrors and interdigitated transducers, and then concatenate them together (with optional free propagation area) to model a full SAW resonator.

¹¹The rules for cascading P-matrix elements are described in Chapter 11 of [160]. Python code that generates P-matrix elements for simple periodic SAW structures and cascades P-matrices automatically may be found in [170].

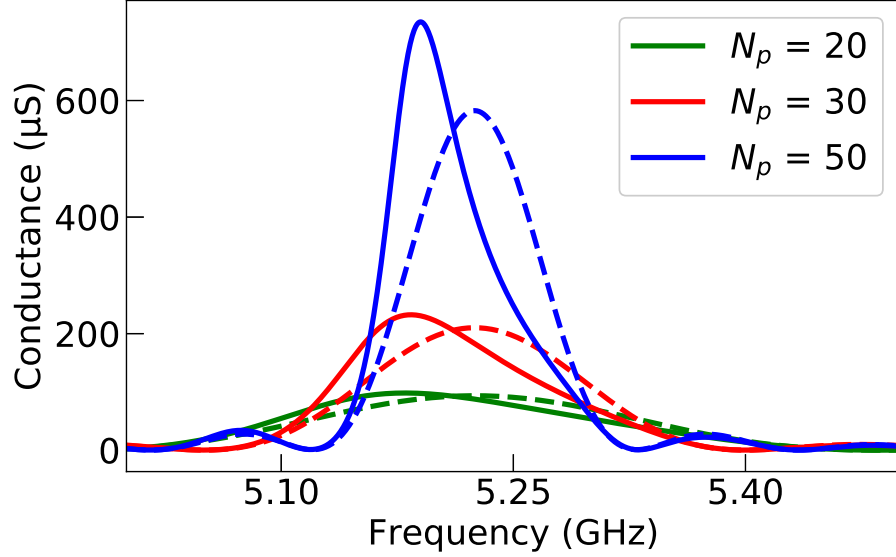


Figure 5.5: **Effect of reflections on IDT conductance:** $\text{Re}(P_{33}(f)) = \text{Re}(Y(f)) = G(f)$ calculated for several values of N_p , with 1% reflection (solid, $r_s = -0.01i$) and no reflection (dashed, $r_s = 0$) per transducer electrode. All $G(f)$ are calculated for a device on ST-X quartz, with a IDT periodicity $\lambda = 600$ nm and finger overlap $w = 60$ μm . When reflection is ignored, the transducers have the same sinc function response as derived in the ideal case in §5.3.1. When we add reflection back in, we see significant distortion of the response as N_p gets larger.

5.3.3 A second look at transducers

To start our analysis of modeling SAW structures with the COM method, we will inspect how reflections modify the response of an IDT compared to the ideal case considered in §5.3.1. Since the main metric of interest is the IDT electrical admittance $Y(f)$, we may make a direct comparison by solving the COM equations for $P_{33}(f)$ with and without reflections taken into account. Doing so for an ideal, reflection-less transducer¹² (i.e. setting the per-grating reflection $|r_s| = 0$) yields the same sinc function response for $\text{Re}(P_{33}(f)) = \text{Re}(Y(f)) = G(f)$ derived in §5.3.1 (Fig. 5.5, dashed lines.) We see the same characteristic dependence of $G(f)$

¹²Even in the presence of reflections, an effectively reflection-less transducer may be created by using a double finger structure, where each finger of the IDT consists of two strips of metal instead of one [167, 11, 17]. However, since this reduces the line-width by a factor of two, at high frequencies these double finger structures become difficult to fabricate.

on N_p , namely that maximum response $G(f_0) \propto N_p$ and the bandwidth of the IDT is $\propto N_p^{-1}$.

When we repeat the same calculation with reflections included ($r_s = -0.01i$, or 1% reflection per IDT finger¹³) we observe significant distortion of the sinc function response. When $N_p \ll |r_s|^{-1}$ we expect reflections to not contribute greatly to the response. Indeed, we see in Fig. 5.5 that for $N_p = 20$ the sinc function response is largely maintained, except for a minor shift of the center frequency. However, as N_p becomes comparable to $|r_s|^{-1}$, the response becomes asymmetric, the bandwidth becomes smaller than in the reflection-less case, and the peak response becomes larger. This means that a device with strong reflection will actually have *stronger* piezoelectric coupling than an otherwise identical non-reflective device, however these devices must be more carefully designed to account for non-idealities like frequency distortion.

5.3.4 Acoustic Bragg mirrors

By taking advantage of the finite reflection per electrode, we may fabricate a SAW *mirror* by placing electrically disconnected strips of metal in the path of SAW propagation. Each strip of metal will reflect some small fraction of an impinging SAW: If the half-wavelength $\lambda_{SAW}/2$ of an impinging SAW is close to the periodicity of these strips, the reflected waves will constructively interfere, satisfying the Bragg condition, and SAWs will be coherently reflected. If the number of strips $N_{\text{mirror}} \gg |r_s|$, the Bragg mirror will reflect impinging SAWs with a high probability, and we may approximate the structure as a high-finesse mirror. The frequency range Δf over which the Bragg condition is met and SAWs are reflected completely in an infinitely long grating (also called the width of the SAW stop-

¹³The phase of the reflection coefficient quantifies what phase the reflected wave will have with respect to the impinging wave, and will in general be imaginary, with the sign depending on the specific device parameters [160].

band), is determined [160] by the magnitude¹⁴ of r : $\Delta f/f \approx 2|r|/\pi$.

The P-matrix reflection coefficient P_{11} calculated by solving the COM equations, and its magnitude and phase are plotted as a function of frequency in Fig. 5.6(a) and (b) respectively. Anticipating assembling mirrors and a transducer into a SAW resonator, for this model we have chosen a $\lambda = 604$ nm grating on ST-X quartz, giving $f_0 \approx 5.19$ GHz slightly lower than the center frequency of the transducers in Fig. 5.5 to account for distortion. Once again, the reflection coefficient per strip is set to $r_s = -0.01i$, i.e. 1% of the SAW amplitude is reflected at each strip. Since the Bragg grating is electrically disconnected, transduction is excluded. We see from Fig. 5.6(a) that in the stop-band, the SAW is reflected with near unit probability, while outside the stop-band the reflection coefficient is a complicated function of frequency.

The phase of the reflected signal evolves smoothly as a function of frequency inside the stop-band: from this, we can calculate the delay time of a SAW reflecting off the structure as $\tau_d(f) = -d\theta_{P_{11}}/df(2\pi)^{-1}$, where $\theta_{P_{11}}$ is the phase of P_{11} . This relationship gives us a simplified model of the Bragg mirror: rather than considering the grating as a distributed mirror, inside the stop band we may think of the Bragg reflector as a point-like mirror positioned at distance $L_P(f) = v\tau_d(f)/2$ from the edge of the grating (see Fig 5.7(a).)

5.3.5 SAW resonators

Having developed models for both IDTs and Bragg mirrors, we are now in a position to model the electrical response of a SAW resonator. A SAW resonator (see Fig. 5.7(a) for a

¹⁴Without doing any calculations, we may gain intuition about this relationship by remembering that in a Bragg mirror, SAWs are reflected over a finite distance which is roughly inversely proportional to the reflection per grating. If the SAWs were completely reflected by a single strip (a point-like mirror) there would be no Bragg condition required for reflection, and SAWs of all frequencies would be reflected. Conversely, as $|r| \rightarrow 0$, the length of a Bragg mirror required to reflect the SAW becomes infinitely long and the range of frequencies Δf that meet the Bragg condition in an infinitely long structure approaches zero.

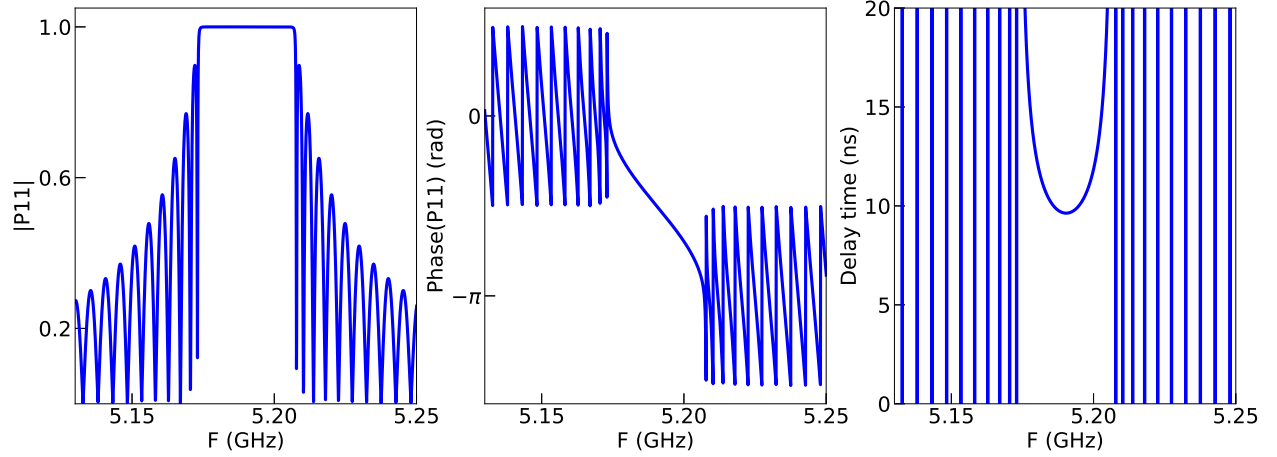


Figure 5.6: **Bragg Mirrors:** (a) Magnitude and (b) phase of the reflection coefficient $P_{11}(f)$ as a function of frequency calculated for a SAW Bragg mirror. This device, calculated for a mirror fabricated on ST-X quartz, has 940 reflector gratings spaced by $\lambda/2 = 302$ nm (total device length ≈ 284 μm) and assumes that each grating reflects 1% of an incoming SAW ($r_s = -0.01i$.) Over the mirror stopband, SAWs are reflected with near unit probability and the phase of the reflection coefficient evolves smoothly as a function of frequency. (c) Delay time of a SAW reflected off the mirror $\tau_d(f) = -d\theta_{P_{11}}/df(2\pi)^{-1}$. From this, we may calculate the effective penetration length of the mirror $L_P(f) = v\tau_d(f)/2$.

schematic) consists of an IDT sandwiched between two Bragg mirrors. Quite like an optical Fabry-Pérot cavity, the resonance condition of a SAW resonator is satisfied when the effective length of the cavity matches a half-integer number of SAW wavelengths. Unlike its optical counterpart, the external coupling to the resonator is provided by the IDT, which inherently lives in the SAW propagation path. While this has the advantage of increasing the coupling strength (since the SAW may interact with the external circuit over many wavelengths), it comes with several caveats, which will be explored further in Chapter 6.

To model a SAW resonator, we will take the two structures we have already modeled (IDTs and Bragg mirrors) and use the cascading properties of the P-matrix to build a model for a SAW resonator. An example calculation is plotted in Fig. 5.7: it assumes SAW propagation on ST-X quartz, launched and transduced by an IDT with periodicity

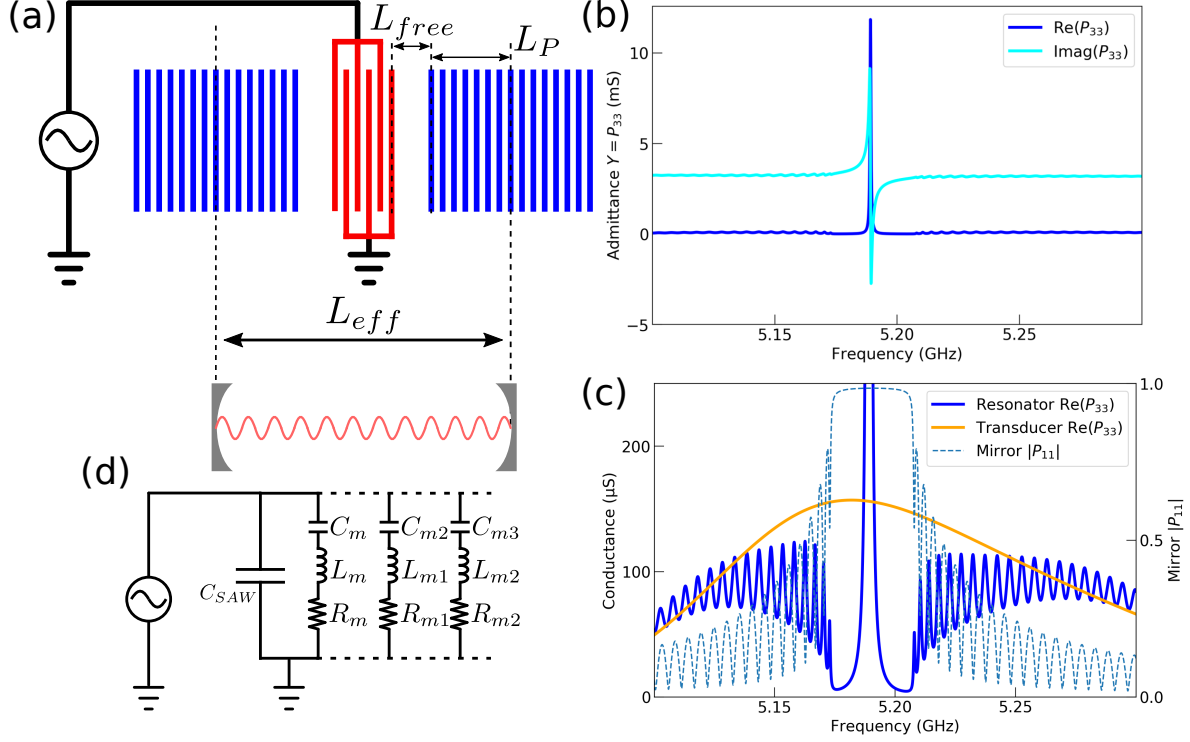


Figure 5.7: **COM model of a SAW resonator:** (a) Schematic of a SAW resonator, consisting of an IDT (red) for external coupling and two Bragg mirrors (blue) on either side of the IDT confining emitted SAWs. This structure may be thought of as an acoustic Fabry-Pérot resonator, with length L_{eff} determined by the IDT width, the mirror spacing L_{free} and the effective penetration length of SAWs into the mirror L_P . (b) Real and imaginary parts of the SAW resonator electrical admittance P_{33} plotted as a function of frequency, modeled by concatenating the P-matrices of each individual structure, see text for details. The response plotted in (b) inherits properties from both the transducer admittance P_{33} and the mirror reflection P_{11} , as is evident from the zoomed in plot of the conductance (c) plotted alongside the mirror P_{11} and transducer conductance $\text{Re}(P_{33})$. However, the full response is well approximated by an effective series LRC resonance in parallel with the geometric IDT capacitance C_{SAW} , the so-called Butterworth-van Dyke equivalent circuit, shown in panel (d). If the free spectral range of the structure is low enough such that more than one resonance condition may fall within the stop-band, we may model the multiple resonant structure by adding more LRC resonances in parallel with the structure.

$\lambda = 600$ nm, width $w = 60$ μm , and $N_p = 25$ finger pairs. The mirrors on either side of the IDT are identical to the structure described in Fig. 5.6. Note that the mirrors have a periodicity $\lambda_{\text{mirror}} = 604$ nm so that the mirror stop-band is centered on the distorted IDT response. There are several important caveats to this model worth mentioning:

1. To model propagation loss, and make the any resonances have finite width, this model includes a small imaginary component to the SAW wave vector.
2. Sometimes, for design purposes or to ease fabrication, it makes sense to add some distance of free SAW propagation L_{free} between the IDT and the Bragg mirrors. To account for this, rather than cascading 3 P-matrixes together as “Mirror-IDT-Mirror”, this model actually consists of 5 P-matrices cascaded together to create a more realistic “Mirror-empty space-IDT-empty space-Mirror” model. In this model, we have chosen $L_{free} = 3 \mu\text{m}$.

In Fig. 5.7(b) we plot both the real and the imaginary parts of the cascaded $P_{33} = Y(f)$. The model resonator displays a clear resonance (zero crossing of $\text{Im}(Y(f))$) along with enhanced $\text{Re}(Y(f))$ at $f \approx 5.19$ GHz. If we zoom in on $\text{Re}(P_{33})$ we see that the fine-scale structure inherits properties of both the Bragg mirrors and the IDT. Inside the stop-band of the mirrors, the conductance is suppressed (except at the resonant frequency, where it is enhanced), while outside the stop-band the conductance is a convolution of the IDT conductance and the mirror reflectance (plotted alongside in orange and dashed blue respectively.) The fine scale features outside the stop band are, however, somewhat inconsequential compared to the response of the structure near its resonance. This resonant response may be accurately modeled as fictitious LRC series circuit in parallel with the geometric capacitance of the IDT, as shown in Fig. 5.7(d). Thus, the underlying piezoelectric interaction of the resonator structure may be reduced to a simple LCR resonator, called its Butterworth-van Dyke equivalent [171]. This approximation greatly simplifies the analysis of integrating SAW devices with other electronic circuits, such as superconducting qubits, which will be discussed further in Chapter 6.

A hallmark of Fabry-Pérot resonators is that they may host multiple modes, at every frequency that satisfies the resonance condition that the cavity length L_{eff} is equal to a half-integer number of wavelengths. For an ideal point-like mirror, which reflects all wavelengths, there are infinitely many resonant frequencies, spaced by the free spectral range $\Delta f_{FSR} = v/2L_{eff}$. However, our Bragg gratings are not ideal point-like mirrors: as we saw in § 5.3.4 they will only efficiently reflect over a small stop-band of width $\Delta f/f \approx 2|r_s|/\pi$. How many resonant conditions do we expect to satisfy inside a cavity? Will these structures even necessarily *have* a resonant frequency?

Since the Bragg gratings reflect roughly $|r_s|$ of the impinging SAW per grating, and the gratings are spaced by $\approx \lambda/2$, a good approximation of the mirror penetration length is $L_p \approx \lambda/2|r_s|$. Therefore, the *minimum* length of a SAW cavity, is $L_{eff} = 2L_p = \lambda/|r_s|$. For such a cavity, $\Delta f_{FSR}/f = v/2L_{eff}f = |r_s|/2$. Since $1/2 < 2/\pi$, Δf_{FSR} will be slightly smaller than the stop-band width, and we expect the cavity to host at least one mode. We also see that as long as the space between the mirrors is smaller than L_p , the cavity will usually host only one mode. Of course, if we add in a large free-propagation distance d , Δf_{FSR} will decrease and the cavity may host many modes [11, 169]. In the case of multiple resonances, we may model each resonant condition as additional parallel LRC resonance (as displayed in Fig. 5.7(d) by the elements connected by dashed lines.)

Chapter 6

Quantum acoustics using surface acoustic waves

The wide applicability of surface acoustic waves to microwave telecom frequency electronics, which superconducting qubits are more or less specifically tailored to be compatible with, has not gone unnoticed. Not long after O’connell et al. demonstrated coherent coupling between a superconducting qubit and a MEMS resonator [2], several groups around the world began reporting results using piezoelectric substrates to mediate coupling between superconducting qubits, mainly transmons, and surface acoustic waves [7, 8, 9, 10, 11, 17, 12, 172, 15]. Indeed, transmon qubits have a major advantage when considering interaction with SAWs: the IDT can play the role of both the SAW transducer and the shunt capacitance that lowers E_C and brings the qubit into the transmon regime.

The slow speed of sound opens up many interesting potential studies/applications of SAWs in the quantum regime that would not necessarily be accessible with light, or even other acoustic modes such as bulk acoustic waves or microelectromechanical devices. Importantly, SAWs propagate on the *surface* of a substrate, where devices are lithographically defined. An IDT structure interacts with a SAW over many wavelengths, meaning that in quantum acoustics with SAWs the “atom” (qubit) is typically much larger than the wavelength of the radiation interacting with it. This is in stark contrast to quantum optics, where the quantum

system (an actual atom or a superconducting qubit) is almost always much smaller than the wavelength of light. This parameter regime has led to the prediction and observation of exotic effects, such as non-exponential decay of a qubit [16, 17] whose decay is dominated by conversion into SAWs. Additionally, since a qubit may interact with SAWs over multiple wavelengths, the acoustic coupling rate to a SAW mode may conceivably be made much larger than its electromagnetic counterpart. It is foreseeable that quantum acoustics devices may be able to enter the so called ultrastrong coupling regime, where the excitation exchange rate Γ between the qubit/SAWs is of order ω_{01} [173], though no experiments exploring the regime have yet been reported.

In classical microwave electronics, the slow speed of SAW propagation is often used to create delay lines, where a SAW signal may spend an appreciable amount of time (several μs) propagating across a substrate before it is transduced again. This pitch-and-catch scheme of a SAW delay line can foreseeably be used as a resource in quantum devices. Several demonstrations of employing this delay have already been experimentally recorded, including using the SAW propagation delay to implement a delayed choice quantum erasure experiment [15] and using a qubit as a “router” to controllably block propagating SAWs [174]. Additionally, the small wavelength of SAWs at $\sim\text{GHz}$ frequencies have been proposed as a resource for creating bosonic quantum memories using SAW resonators [18], since a resonator with a high density of acoustic modes can easily be fabricated with a relatively small spatial footprint. Quantum acoustics devices have also been proposed for use as coherent microwave-to-optical converters [175], where the wavelength of the SAW matches the wavelength of a confined optical mode.

This field of quantum acoustics, using SAWs or other modes, is in its infancy, with much uncharted territory to be explored both answering fundamental scientific questions and

designing potential applications. When thinking about how to start exploring this territory, we would like to ask “how can we leverage the experimental apparatus built up earlier in this thesis (i.e. 3D transmon qubits) to start performing quantum acoustics experiments?” In this chapter, we outline experiments exploring a novel way of mediating the SAW-qubit coupling in the 3D transmon geometry: a capacitively coupled device in a “flip-chip” configuration. We will explore how one may model such an experiment, and then report on preliminary experimental data investigating these devices.

6.1 Mediating the coupling between SAWs and qubits

The first question we may ask is “how can we mediate coupling between SAWs and qubits?” In light of the previous section, this seems like a somewhat unnecessary question: we already established that transmon qubits need a shunt capacitance, which can easily be provided by the IDT. Therefore the most conceptually straightforward method seems to be direct galvanic coupling, where either side of the Josephson junction is directly connected to either electrode of the IDT. This design is conceptually straightforward, has been pursued by multiple groups [7, 11, 9], and in principle offers the strongest coupling, since voltage fluctuations across the IDTs induce voltage fluctuations across the junction with unit efficiency.

There are, however several drawbacks to this direct galvanic coupling scheme. For instance, the anharmonicity of the transmon is inversely proportional to the IDT capacitance, and thus there is a limit to the size of the transducer (both finger pair number N_p and overlap length w) before the anharmonicity gets too small¹. Since, as we saw in Chapter 5, the

¹Here, “too small” is set by the control pulses we send to the device: we would like to remain in the qubit manifold (i.e. not accidentally excite the transmon to its $|2\rangle$ state) while using pulses of finite width. These pulses are sinusoidal, however their finite width means they will inevitably have Fourier components at higher frequencies. We may accidentally drive the $|1\rangle \rightarrow |2\rangle$ transition if there are significant Fourier components at frequencies $\geq \alpha$. Since control pulses tend to be ~ 10 ns long, we would at least like $\alpha > 100$ MHz.

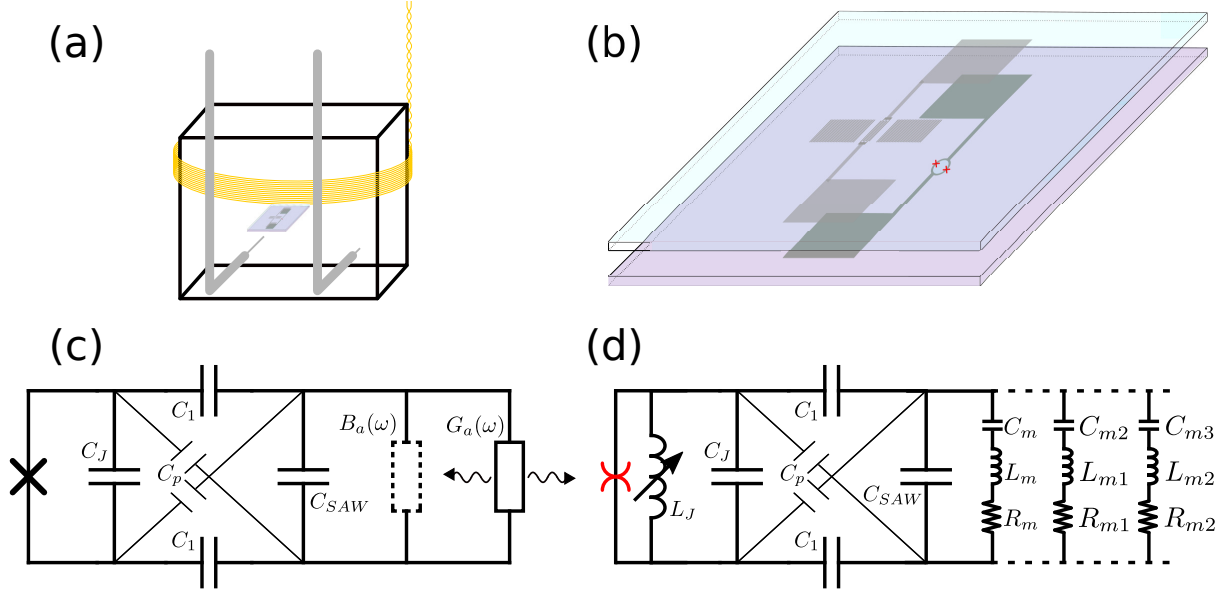


Figure 6.1: **Schematic of the proposed experiment:** The proposed quantum acoustics experiment consists of a SAW resonator capacitively coupled to a superconducting qubit device in a flip-chip configuration (b), read out using a 3D electromagnetic microwave resonator (a). (c) The device may be modeled as a superconducting qubit coupled to a SAW admittance via a capacitive network. (d) We may further approximate the SAW by it's Butterworth-van Dyke equivalent circuit, and model the transmon qubit as a weak nonlinearity in parallel with a linear circuit to extract the Hamiltonian of the system.

acoustic conductivity of a SAW IDT is proportional to w and N_P^2 , this puts a limit to the achievable coupling in practice. Additionally, Josephson junctions are notoriously sensitive to static discharge, and strongly piezoelectric substrates like LiNbO_3 are often susceptible to discharging events² and thus unwieldy to fabricate Josephson junctions on. It is, therefore, desirable to fabricate the qubit device on a well-known and well-behaved substrate like silicon, and to couple it to a SAW device on another piezoelectric substrate.

Along these lines, there have been demonstrations of coupling a qubit to a SAW device on another substrate using a tunable mutual inductance [10, 172, 15]. While these devices have demonstrated high-coherence and flexibility, the experimental setup is somewhat involved,

²Lithium niobate is also pyroelectric, i.e. temperature changes can induce polarization of the material. My experience has been that, when heating devices to cure resist in fabrication, it is possible to hear static discharge events in the substrate!

and incompatible with the 3D transmon geometry, where we aren't necessarily able to run bias lines into the cavity to tune individual circuit elements, and the qubit device is explicitly disconnected from ground. A more simple way to engineer coupling between the SAW and qubit devices while still fabricating them on different substrates is through a geometric capacitance, outlined in Fig. 6.1. This coupling is similar to the coupling between a transmon and a linear LC oscillator described in Chapter 2: a SAW device, which we choose to be an acoustic resonator in analogy to the electromagnetic resonator of cQED, fabricated on a piezoelectric substrate is coupled to a Josephson junction device on a silicon substrate. We then place one device face down on the other: in this “flip-chip” configuration, the two sets of pads serve a dual purpose: they form the parallel plate capacitance C_1 that couples the SAW device to the qubit, as well as the dipole antenna that couples the qubit to the 3D electromagnetic cavity we use for manipulation/readout.

6.2 Device design

To understand the design constraints of this experiment, we will need to consider a Hamiltonian that describes the coupled qubit-acoustic system. For a first experiment, we would like to couple a discrete acoustic mode, confined in a SAW resonator, to a superconducting qubit. This setup is an acoustic analogue to cQED, and therefore it makes sense to borrow from the formalism developed in Chapter 2 where a qubit was coupled to a discrete electromagnetic mode. Therefore, to model this experiment, we choose the by now familiar Jaynes-Cummings-like Hamiltonian [18]

$$\mathbf{H} = \hbar\omega_{01}\mathbf{q}^\dagger\mathbf{q} - \frac{\alpha}{2}\mathbf{q}^\dagger\mathbf{q}^\dagger\mathbf{q}\mathbf{q} + \sum_k \hbar\left[\omega_k\mathbf{m}_k^\dagger\mathbf{m}_k + g_k(\mathbf{q}^\dagger\mathbf{m}_k + \mathbf{q}\mathbf{m}_k^\dagger)\right] \quad (6.1)$$

Here $\mathbf{q}^\dagger$ and $\mathbf{q}$ are the raising/lowering operators of the qubit mode, ω_{01} is the ground to excited state transition frequency of the transmon qubit, $\alpha = E_C$ is the quartic nonlinearity of the transmon circuit. When designing the experiment, we opt to explicitly maintain this nonlinear term rather than truncating to the qubit manifold and writing down the Hamiltonian in terms of Pauli operators, since the transmon capacitance that determines E_C will be a major design parameter for determining if we have enough nonlinearity to operate the circuit in the transmon regime.

As noted in Chapter 5, a single acoustic resonator may host multiple modes. Thus, the Hamiltonian Eqn. 6.1 consists of a sum over k acoustic modes, with raising and lowering operators $\mathbf{m}_k^\dagger$ and $\mathbf{m}_k$, each a linear oscillator oscillating at ω_k with a coupling g_k to the transmon mode. In this dissertation, however, we will only be considering geometries with a single acoustic mode: in this case, the Hamiltonian reduces down to

$$\mathbf{H} = \hbar\omega_{01}\mathbf{q}^\dagger\mathbf{q} - \frac{\alpha}{2}\mathbf{q}^\dagger\mathbf{q}^\dagger\mathbf{q}\mathbf{q} + \hbar\omega_m\mathbf{m}^\dagger\mathbf{m} + \hbar g_m(\mathbf{q}^\dagger\mathbf{m} + \mathbf{q}\mathbf{m}^\dagger) \quad (6.2)$$

Of the four parameters in this Hamiltonian, the SAW resonance is relatively easily constrained by λ_{IDT} , and ω_{01} may be tuned *in situ* by using a split-junction transmon and applying a magnetic flux. Thus, the main design parameters we must extract from any proposed experimental geometry are the transmon anharmonicity $\alpha = E_C = e^2/(2C_{eff})$, where C_{eff} is the effective capacitance of the transmon mode, and the qubit-acoustic coupling g_m . The modeling process described below aims to calculate C_{eff} and g_m from a given geometry. To facilitate this discussion, we will consider an example device throughout this section, with design parameters and derived values summarized in Table 6.1.

6.2.1 Modeling the experiment

As a starting point, let us take the effective circuit diagram of the device outlined in Fig. 6.1(c). The qubit device on silicon consists of a Josephson junction in parallel with a capacitance C_J that incorporates both the intrinsic junction capacitance and the capacitance between the paddles. Each paddle forms a parallel plate capacitor C_1 with the matching paddles from the SAW device: in the symmetric case, these parallel plate capacitances are equal. The SAW IDTs have some intrinsic capacitance as well as the capacitance of the paddle structures connected to the IDT: these add in parallel to the total capacitance between the two electrodes of the SAW device C_{SAW} . As explained in Chapter 5, the electrical response of the SAW cavity is encapsulated in a frequency dependent admittance $Y_{SAW}(f) = G_a(f) + iB_a(f)$. There will also be some parasitic capacitance C_p between the non-parallel plate pads, which we would like to minimize. We justify this lumped element circuit model by noting that electromagnetic radiation < 5 GHz has a wavelength of 6 cm, and the SAW+qubit device is ≈ 1 mm in its longest dimension.

The comparatively weak nonlinearity of the transmon, such that the Hamiltonian can approximately be diagonalized in the harmonic oscillator basis, offers a distinct advantage when designing experiments that involve transmon qubits: when engineering a Hamiltonian, it is often fairly accurate to break up the transmon into an effective linear inductor in series with a small, nonlinear element. In this manner, we may model the circuit as a weak nonlinearity connected to a linear network of capacitors and inductors, which form a system of *linear* oscillators, and consider the nonlinearity as a perturbation on this system [78]. We represent this schematically by splitting the Junction into an inductor $L_J = \Phi_0^2/E_J$ (where E_J is tunable by an external magnetic flux) along with a parallel nonlinearity (red element

in Fig. 6.1(d).) We then quantize the Hamiltonian of the system in the basis of (coupled) harmonic oscillators.

It was belabored in Chapter 2 that a quantum harmonic oscillator, in the absence of any anharmonicity, responds with a spectrum equivalent to its classical counterpart. This is also true for a set of *coupled* harmonic oscillators [62], i.e. a network of inductors and capacitors. Thus, the problem of reconstructing the Hamiltonian of a device for a given design is greatly simplified: all we must do is simulate the *classical* response of the circuit (i.e. calculate $Y(f)$ of the linear network), and find the resonances (where $\text{Im}(Y(f)) = 0$) and characteristic impedances of these resonances. Equipped with this knowledge, we can promote these classical oscillators to a set of coupled quantum harmonic oscillators, with zero point fluctuations described by Eqn. 2.24, and add the transmon nonlinearity back in as a perturbative term.

Pragmatically we will break this up into several steps: we will use a commercial finite element solver to numerically solve for the capacitances between the macroscopic electrodes for a given device geometry. We will then use the coupling of modes model described in Chapter 5 to construct the Butterworth-van Dyke (BvD) equivalent of the SAW resonator. This will numerically fill out the circuit diagram shown in Fig. 6.1(d): we will then solve this circuit for $Y(f)$ as a function of the tunable linear Josephson inductance L_J . From this data, we will be able to extract the coupling between the linear modes g_m and the effective capacitance of the transmon-like mode, which will give us α .

6.2.1.1 Finite element modeling

For maximal coupling between the SAW resonator and the superconducting qubit device, we would like to maximize the capacitance C_1 and minimize the parasitic capacitance C_p , see

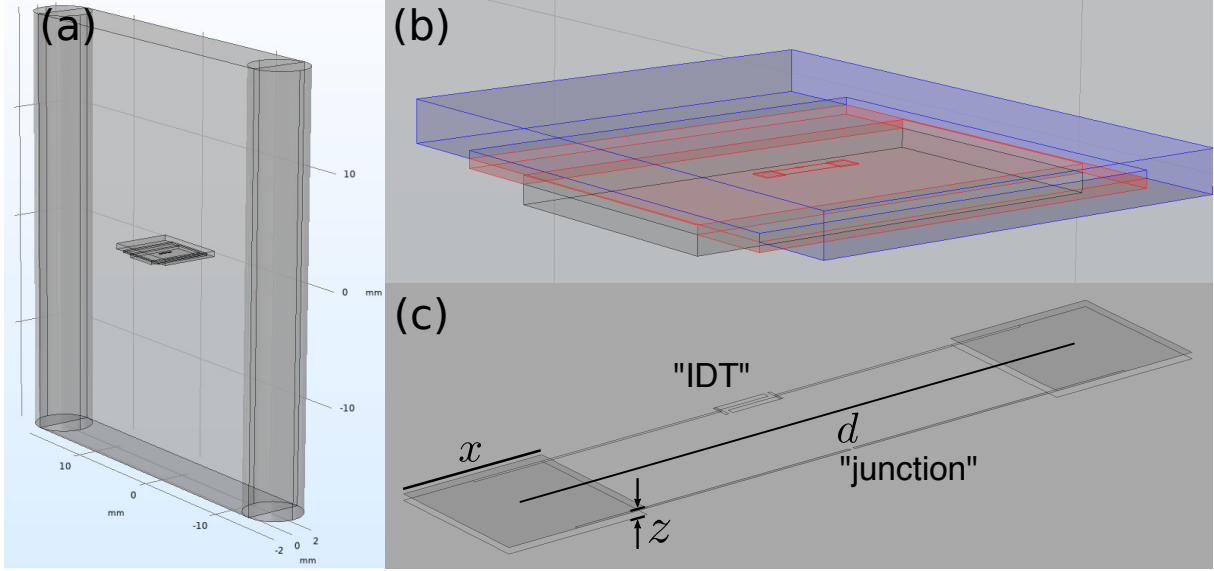


Figure 6.2: **Finite element simulations of the paddle structure:** We use a commercial finite element solver (COMSOL) to solve for the capacitance network between each node in the device circuit. (a) Ground potential is set at the walls of the cavity, while inside the cavity we simulate the presence of dielectrics (b) corresponding to the SAW substrate (blue) and the silicon wafer grey box.) Close to the device, we define a volume (red) with increased mesh resolution to increase accuracy/efficiency of the simulation. (c) The device consists of four conducting “ports” (nodes), two on the surface of each substrate. Each node has a square pad with linear dimension x , and a $5\text{ }\mu\text{m}$ wide lead connecting the pad to the device (IDT or junction.) The centroids of the pads are separated by d , and the spacing between the substrates is z .

Fig. 6.1(c). We would also like the device to have strong coupling to the electromagnetic 3D cavity, so we can perform control and readout on the transmon qubit degree of freedom. To satisfy these constraints, we choose a geometry outlined in Fig. 6.2(d): both the qubit and the SAW device consist of square pads with linear dimension x , whose centroids are spaced by d , connected to the active element by a $5\text{ }\mu\text{m}$ wide lead. These pads, when separated by a distance $z \ll x$ away from each other in a flip-chip configuration, form a parallel plate capacitance that defines C_1 . The leads that lead to the transmon/SAW devices at the center of the device are thin, and placed on opposite sides of the device to minimize C_p .

To obtain accurate values for the capacitive network, the grounding of the electric fields

must be accurate: when one electrode of the device is raised in voltage relative to the other electrodes, ground is not at “infinity”, but rather at the walls of the cavity, which are somewhat close to the device relative to its size. Therefore, we encapsulate the simulation with a conducting boundary condition having the same geometry as the inside of the electromagnetic cavity (see Fig. 6.2(a).) Additionally, we must define volumes to model the dielectric host substrates, see Fig. 6.2(b). This is fairly straightforward for materials like silicon and quartz, which have a single dielectric constant, however LiNbO₃ has an anisotropic dielectric tensor, and so the particular orientation of the SAW device with respect to the underlying crystal lattice matters for modeling the capacitance³.

The feature sizes of the device are small (several μm) while the overall volume of the simulation is large, of order $\sim\text{cm}^3$. Given this disparity, it is beneficial to split the element mesh into two parts: one part close to the device (red volume in Fig. 6.2(b)) with an extremely fine mesh, and the rest of the volume with a more coarse mesh. This allows for accurate simulation of the smaller structures without an overwhelming number of elements. Since the IDT’s are numerous, and have linear dimensions in the ~ 100 nm range, it would be unwieldy to simulate them even using this split-mesh approach: we therefore simulate a dummy structure in place of the IDT’s, and add the analytic IDT capacitance C_{IDT} to the total capacitance of the SAW structure C_{SAW} by hand. For the purposes of this simulation, we also ignore the small geometric capacitance of the Josephson junction.

From our example device summarized in Table 6.1, we have pad size $x = 250 \mu\text{m}$, pad spacing $d = 1 \text{ mm}$, and wafer separation $z = 10 \mu\text{m}$. Using these values, we find

³A summary of physical properties of LiNbO₃, including the values of the dielectric tensor, may be found in [176], and the Euler angles to rotate axes to common SAW cuts of LiNbO₃ may be found in Chapter 4 of [160]. Note that at microwave frequencies, one should use the values of the dielectric tensor at constant strain ϵ_{ij}^S .

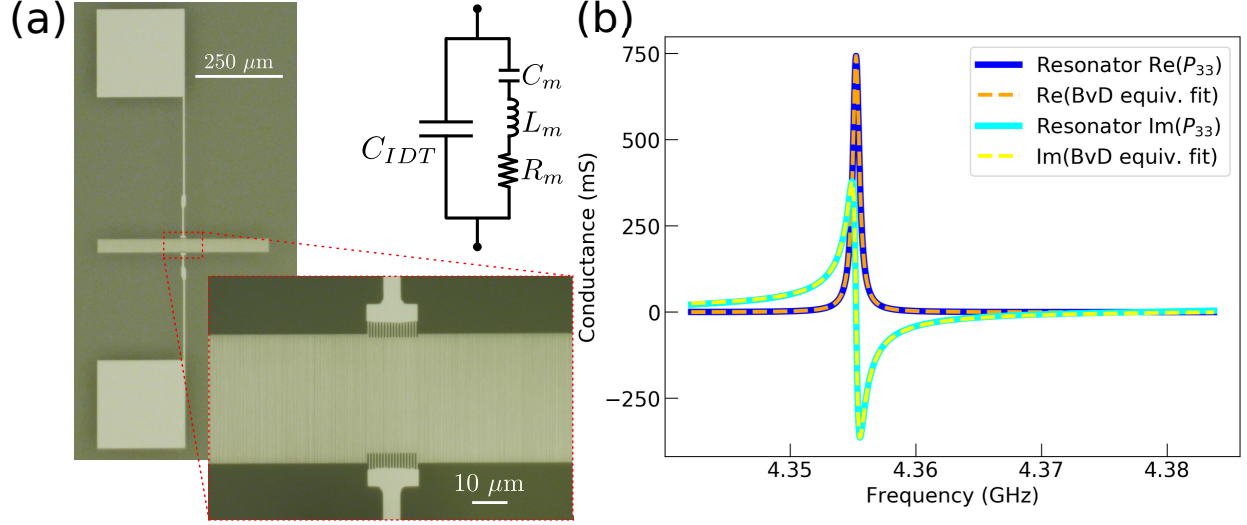


Figure 6.3: **Butterworth-van Dyke equivalent response of a SAW resonator:** (a) Microscope image of an actual SAW device fabricated on LiNbO_3 , with design parameters corresponding to those listed in Table 6.1. The electrical response of the resonator may be accurately modeled by the Butterworth-van Dyke (BvD) equivalent circuit: the IDT geometric capacitance C_{IDT} in parallel with a fictitious LCR resonance. (b) Real and imaginary parts of $Y(f) = P_{33}(f)$ of the SAW structure depicted in (a), calculated using the coupling of modes method. Dashed lines: a fit to the admittance of the BvD equivalent circuit, showing a high degree of accuracy.

$C_1 = 117.8$ fF, $C_p = 34.3$ fF, $C_J = 21.3$ fF, and the portion of the SAW capacitance associated with the pads $C_{SAW} - C_{IDT} = 102.6$ fF. The large anisotropy between C_J and the pad proportion of C_{SAW} may be attributed to extremely large values of the dielectric tensor of LiNbO_3 relative to the dielectric constant of silicon.

6.2.1.2 Extracting the Butterworth-van Dyke equivalent circuit

Having extracted the coupling capacitances, we now need to model the SAW resonator to extract the BvD equivalent circuit. We use the coupling of modes (COM) model, as outlined in Chapter 5, to generate a P-matrix for each SAW structure and cascade single structure P-matrices to create a model for the entire SAW resonator. We then fit $P_{33}(f) = Y(f)$ and extract from the fit BvD equivalent circuit elements that represent the electrical response of

the SAW device accurately.

The example resonator from Table 6.1, which has been fabricated and is pictured in Fig. 6.3 has an IDT with $\lambda_{IDT} = 900$ nm and $N_p = 16$ finger pairs that have $w = 35$ μ m overlap. There are 525 Bragg gratings on either side of the IDT, with periodicity $\lambda_{mirror}/2 = 456$ nm and no free propagation distance between the IDT and the gratings ($L_{free} = 0$.) We model this structure using a Python based COM modeling package (code may be found in [170]), and inputting a guess for the reflection per unit length ($|r_s| = 1.5\%$, in line with other literature in the field [10]) and a guess for the propagation loss. Using the process outlined in Chapter 5, we generate a cascaded P-matrix for the device given these parameters: the real and imaginary parts of $P_{33}(f) = Y(f)$ inside the mirror stop band are plotted in Fig. 6.3(b).

The main claim of the BvD model is that this structure should be accurately represented by a series LRC circuit in parallel with the geometric capacitance C_{IDT} . The admittance across such a structure, shown in Fig. 6.3(a), is given by

$$Y(\omega) = \frac{R_m(\omega C_m)^2}{(R_m C_m \omega)^2 + (1 - L_m C_m \omega^2)^2} + i \left(\omega C_{IDT} + \frac{C_m \omega - C_m^2 L_m \omega^3}{(C_m R_m \omega)^2 + (1 - C_m L_m \omega^2)^2} \right) \quad (6.3)$$

where $\omega = 2\pi f$. We fit the $\text{Re}(Y(f))$ to this function, yielding R_m , C_m , and L_m as fit parameters. For the model of the SAW resonator described above, $R_m = 1.34$ Ω , $L_m = 0.338$ μ H, and $C_m = 3.95$ aF. The fit to both the real and imaginary parts of Eqn. 6.3 are plotted as dashed lines in Fig. 6.3(b), and agree with the numerically derived $Y(f)$ to a high degree of accuracy within the stop band of the Bragg mirrors.

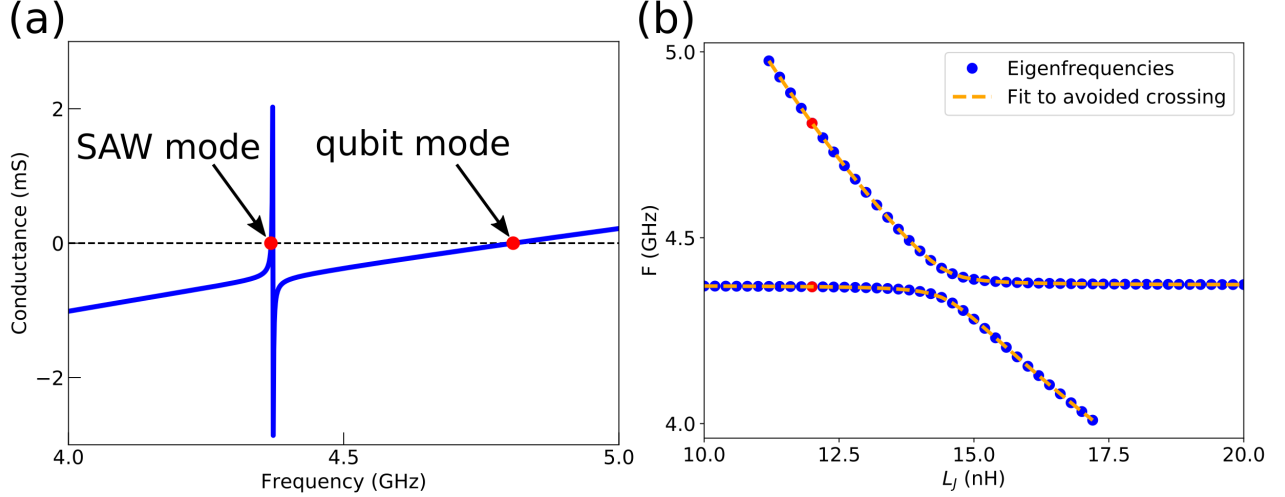


Figure 6.4: **Extracting the Hamiltonian parameters from classical circuit simulations:** (a) $\text{Im}(Y(f))$ of the linear component of the circuit in Fig. 6.1(d), using element values extracted from FEM/COM modeling and $L_J = 12$ nH. The zero crossings (red dots) signify a resonance of the circuit. (b) For each value of L_J , we extract the zero crossings of $\text{Im}(Y(f))$ (blue dots, red dots correspond to data from panel (a)) and fit to extract the Hamiltonian parameters.

6.2.1.3 Classical circuit simulations

We are now equipped with the circuit elements that make up the model of our device outlined in Fig. 6.1(d). The response of the linear segment of the circuit may be calculated using fundamental circuit analysis techniques, and reduced to a single admittance $Y(f)$ in parallel with the non-linear component. In Fig. 6.4(a), we plot $\text{Im}(Y(f))$ as a function of frequency using the element values obtained from finite element simulations of the capacitive structure/COM modeling of the SAW structure. For this plot, we choose $L_J = 12$ nH, which corresponds to a Josephson energy $E_J/h = 13.6$ GHz that is easily obtainable experimentally using a flux tunable transmon. The zero crossings of $\text{Im}(Y(f))$, marked by red dots in Fig. 6.4(a), signify a resonance of the circuit. We identify two modes of the circuit: the SAW-like mode, with a resonant frequency approximately equal to that of the bare SAW resonator $\omega_m/(2\pi)$, and a transmon-like mode far detuned from the SAW mode.

To extract the Hamiltonian parameters, we run the same simulation using different values of L_J , and extract the zero crossings to build the spectrum of the linearized circuit as a function of L_J , plotted in Fig. 6.4(b). We see that the SAW-like mode and the transmon-like mode behave differently as a function of L_J : the SAW mode is more or less constant in L_J , while the frequency of the transmon-like mode shifts depends strongly on L_J . In the uncoupled case, the SAW mode should be at $\omega_m/(2\pi)$ for all values of L_J , while the linearized transmon mode should obey $\omega_q/(2\pi) = 1/\sqrt{C_{eff}L_J}$, where C_{eff} is the effective capacitance of the transmon mode. These modes, however, are coupled and undergo an avoided crossing⁴, shifting the frequencies of the coupled system to

$$f_{\pm}(2\pi) = \frac{\omega_m + \omega_q}{2} \pm \sqrt{g_m^2 + (\Delta/2)^2} \quad (6.4)$$

where, similar to the discussion in Chapter 2, $\Delta = \omega_m - \omega_q$ and g_m is the SAW-qubit coupling we desire. In order to extract g_m and C_{eff} , which gives us $\alpha = E_C = e^2/2C_{eff}$, we fit a single band of the avoided crossing to one of the frequency bands of Eqn. 6.4. A fit to the data is plotted in Fig. 6.4(b) as a dashed orange line, and is shown to match the extracted frequencies well. From this fit, we extract that, for this geometry, we should have $g_m/(2\pi) = 40.2$ MHz, strong enough coupling to be readily observable in an experiment, and $C_{eff} = 93.6$ fF, large enough to put us in the transmon regime but small enough to have sufficient nonlinearity to operate the transmon mode as a qubit.

Design Parameters		Phenomenological Parameters	
Substrate	128°Y-X LiNbO ₃	r_s	$-0.015i$
d	1 mm	Propagation loss	500 Np/m
x	250 μm	Derived Parameters	
z	10 μm	C_1	117.7 fF
N_p	16	C_J	21.3 fF
λ_{IDT}	900 nm	C_p	34.3 fF
w	35 μm	C_{SAW}	383 fF
λ_{mirror}	912 nm	f_0	4.371 GHz
N_{grat}	525	C_{eff}	93.6 fF
L_{free}	0	$g_m/(2\pi)$	40.2 MHz

Table 6.1: Summary of the design parameters, phenomenological parameters, and derived parameters used in the model of the example SAW resonator considered in § 6.2. All design parameters are defined in fabrication.

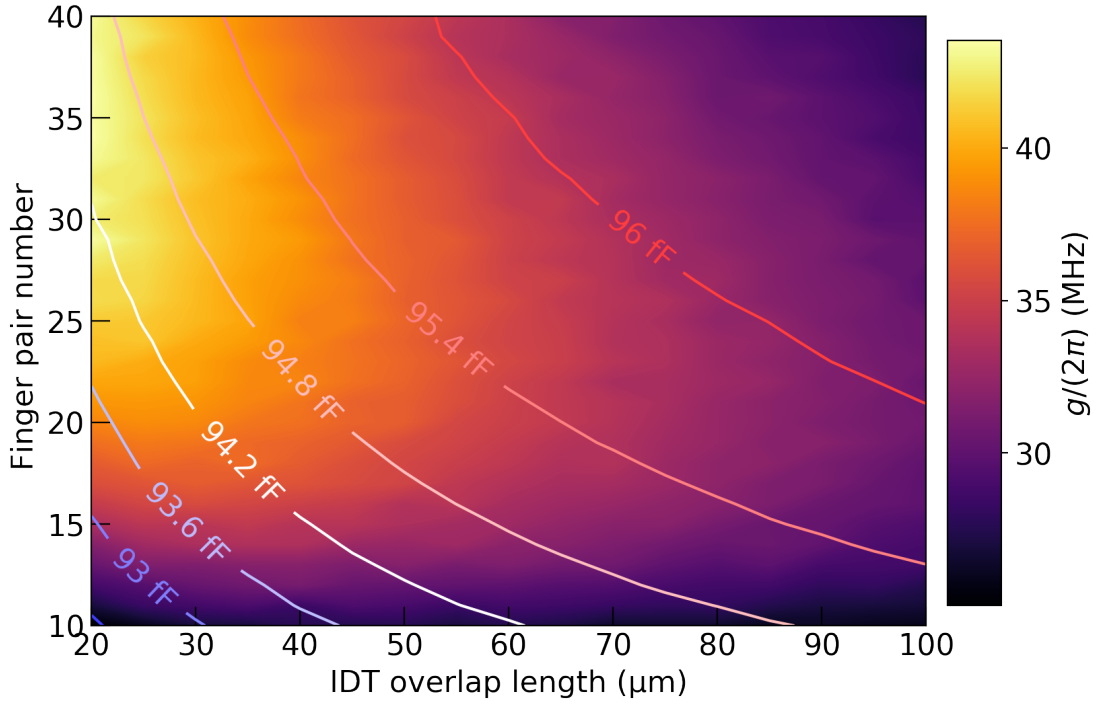


Figure 6.5: **Finite element simulations of the paddle structure:** Color plot: SAW-qubit coupling g_m as a function of finger pair number N_p and IDT overlap length w . Contours: Transmon mode capacitance C_{eff} as a function of finger pair number N_p and IDT overlap length w

6.2.2 Optimizing the SAW device parameters

We are now in a position to try and optimize the experimental geometry to maximize coupling between the SAW resonator and the qubit while maintaining a transmon mode capacitance that we can work with. For the SAW device, aside from λ_{SAW} which sets the operating frequency f_0 , the most important parameters we have control over are the number of finger pairs N_p and the overlap length of the fingers w . How do g_m and C_{eff} depend on these quantities?

In Fig. 6.5, we plot both g_m and C_{eff} as a function of N_p and w , for a device on 128° Y-X LiNbO₃, with $\lambda = 900$ nm and a pad geometry that give the capacitance values outlined at the end of § 6.2.1.1. The first thing to note is that C_{eff} is only weakly dependent upon the geometry of the SAW structure: this makes intuitive sense, since the effective capacitance of several capacitors in series will be dominated by the smallest capacitance. Thus, C_{eff} is largely constrained by the macroscopic pad geometry, and changes in $C_{SAW} \gg C_1$ will have little effect on the overall capacitance of the transmon mode.

Changes in C_{SAW} do, however, strongly effect the transmon-SAW coupling g_m . Inspecting Fig. 6.5, we observe that in the capacitively coupled configuration, g_m monotonically decreases as a function of w . This may be understood along similar lines as how changes in the capacitance in the presence of superfluid helium modified g_0 in Chapter 4. As we increase C_{SAW} , voltage fluctuations across C_{SAW} induce voltage fluctuations across the junction with less efficiency. Thus, increasing w , which from Eqn. 5.32 naively increases the SAW conductance of a bare IDT actually *decreases* the SAW-qubit coupling. Increasing N_p also increases C_{SAW} , however from Eqn. 5.32 we know the bare conductance is proportional

⁴While more often discussed in the context of quantum mechanics, avoided crossings often appear in completely classical systems. A system of any two coupled classical oscillators will undergo an avoided crossing if the resonant frequency of one is tuned through the other.

to N_p^2 . This square dependence wins out, and g_m increases slowly as a function of N_p .

Thus, it seems that to maximize coupling, we should want to maximize N_p while minimizing w . There are, however, some considerations that prevent us from doing this indefinitely. The first is that SAWs emitted from a finite length transducer are not precisely a plane-wave: they undergo some diffraction, and can leak out into the continuum surface unprotected by the Bragg mirrors/unable to be transduced by the IDTs. For short transducers (small w), this leakage ultimately limits the quality factor of the SAW resonator. The diffraction limit of Q is predicted to be $\propto (w/\lambda)^2$ with a constant of proportionality that depends on material parameters [167, 177]. For a first pass experiment, we would simply like the diffraction loss to be small enough that we don't need to worry about it being a dominant decay mechanism: from this perspective, a value $(w/\lambda) \gtrsim 40$ seems reasonable.

Additionally, distortions from intra-IDT reflections may adversely effect optimal coupling. While this can presumably be taken into account using the COM theory as was done in Chapter 5, the exact magnitude of the reflection coefficient r_s for a given set of fabrication parameters seems to be determined empirically in the literature and is not known *a priori*. Thus, it would behoove us to be in a regime where the intra-IDT reflections produce minimal distortion for a wide range of $|r_s|$.

6.3 Device fabrication and assembly

In order to have a predictable acoustic coupling, we must be able to align the paddles that form the parallel plate capacitors with a fair amount of precision. To do so, we employ the mask aligner in the class 100 clean room of the Keck Microfabrication Facility in BPS, which is normally used for aligning substrates to photomasks with high precision for pho-

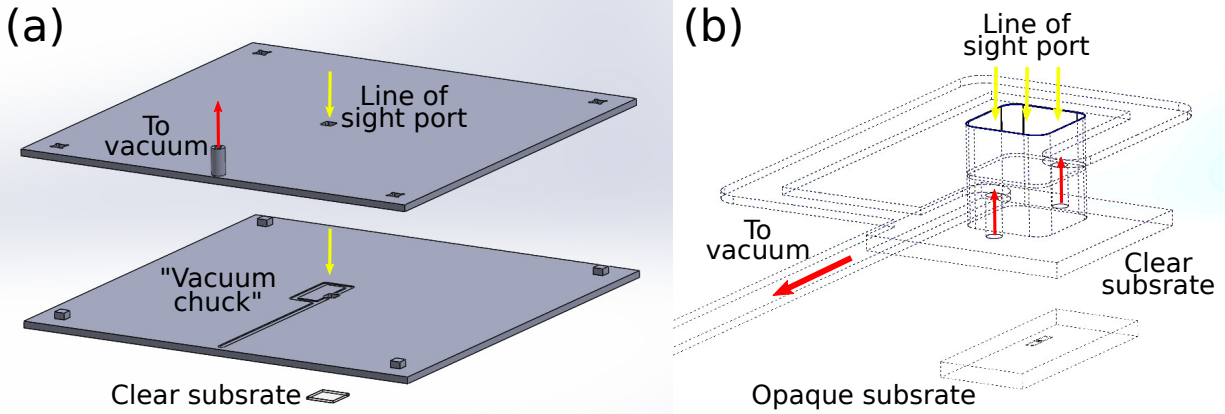


Figure 6.6: **Mask for aligning the flip-chip device stack:** (a) Exploded view of the flip-chip alignment mask. When the two plates are flush, they enclose a small trench that connects two small holes near the center to a hose stub. We tee off the substrate vacuum line to pump on this volume, which provides enough suction to lift a substrate for alignment (red arrows) (b) Zoom in on the center portion of the mask. If we carry a clear substrate with the vacuum chuck, we may use the line of sight port, along with the microscope on the mask aligner, to align the two substrates relative to each other with high precision.

tolithography. In place of a photomask, we use a custom made “mask” designed to fit on the 4” photomask plate on the mask aligner, which is actually two steel plates that fit flush together, see Fig. 6.6. When flush, they enclose a small volume connecting a “vacuum chuck” (two small holes through the bottom mask) to a hose stub, which fits a hose with the same diameter as the substrate vacuum hose on the mask aligner. We tee off the substrate vacuum line to this stub, which provides a strong enough vacuum to pick up a substrate. If the substrate is clear, we can use the line of sight port through the mask to align a device on the top (clear) substrate with a device on another substrate under it, which may be opaque. This is easily done on the mask aligner, which has a built in microscope and 4-axis positioning of the opaque substrate relative to the clear substrate which is held static on the mask.

To control the vertical distance z between the substrates, we use a spacer fabricated from hard baked photoresist, as depicted in Fig. 6.7. Hard baked photoresist has several advantages: photoresist is spun on with a reliable and predictable thickness, and we may

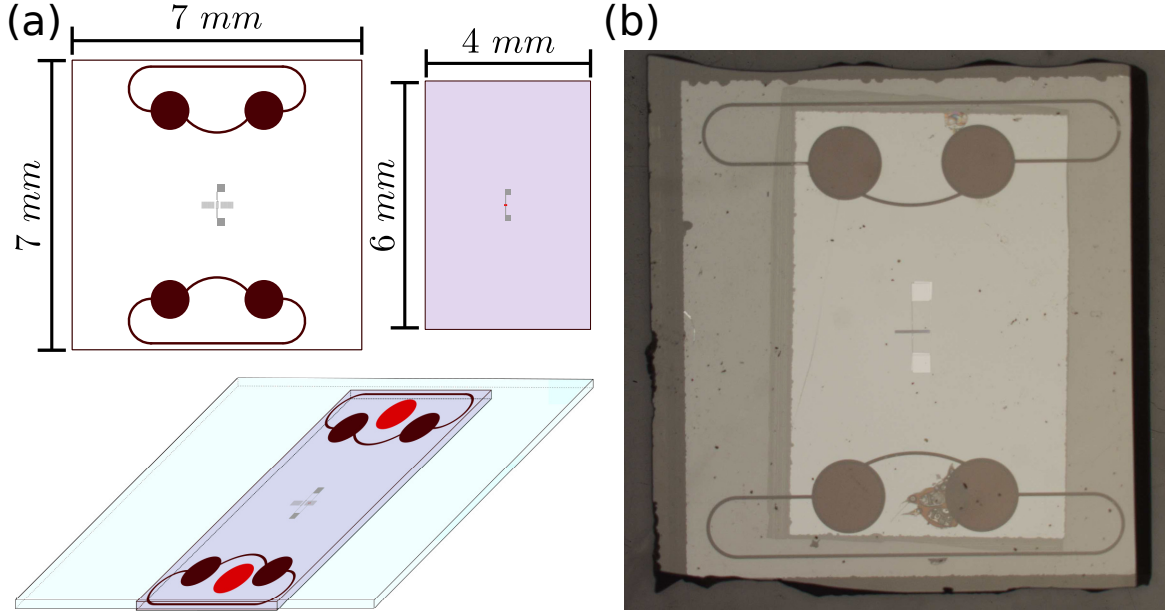


Figure 6.7: **Assembling the flip-chip stack:** (a) Schematics of the devices cut to size, as well as a rendering of the assembled flip-chip stack. The SAW device is made larger than the qubit device for ease of assembly/handling. We use hard baked photoresist (maroon) spacers at the end of the chip to control the spacing between chips, and use soft-baked photoresist (red) to glue the chips together. (b) Image of the flip-chip stack, as seen through the clear LiNbO_3 substrate.

apply multiple layers to generate a thicker spacer. Additionally, once hard baked, photoresist is robust to normal substrate cleaning procedures, which allows us to clean surface after fabricating the spacers. We have fabricated spacer of up to $4\ \mu\text{m}$, though a thicker layer is likely possible. A detailed fabrication recipe may be found in Appendix C. We have observed that the spacer application process destroys the Josephson junction device, and therefore we fabricate the spacers on the piezoelectric substrate. We have also recently observed that, while baking the resist spacers, SAW devices fabricated on 128° Y-X LiNbO_3 almost universally get destroyed by static discharging events. We have, however, been able to successfully fabricate resist spacers on both ST-X quartz and Y-Z LiNbO_3 .

To glue the chips together, we once again employ photoresist [178]: two small dabs of photoresist are placed on one of the devices (generally the qubit device) such that they will

land inside of the spacers once the chips brought together, as shown in Fig. 6.7(a). We then use the mask aligner to align the substrates relative to each other, and bring the chips into contact until the weight of the alignment mask is supported by the chip stack. The photoresist dabs act as glue to hold the chips together, and the weight of the mask ensures that the resist hardens with the chips as close together as possible. We allow the flip-chip stack to sit in place for ~ 10 minutes so the resist can harden slightly, and then transfer it to a hotplate at 110°C , and let it bake for 3-5 minutes to cure the resist such that the stack is fairly robust.

6.4 Experimental results

This section details some preliminary experimental results testing the devices we describe using the model laid out in § 6.2 and fabricate and assemble using the procedure laid out in § 6.3. At this point, several experimental runs have been performed, using both ST-X quartz and LiNbO_3 as piezoelectric substrates. We will focus on a single, exemplary experimental run for each material, and comment on the findings/shortcomings of each.

6.4.1 Devices on ST-X quartz

The first generation of quantum acoustics devices tested were flip-chip stacks with the qubit device on silicon and the SAW device on ST-X quartz. While quartz is a fairly weak piezoelectric material, it has a fairly low dielectric constant ($\epsilon \approx 3.8$) and thus the capacitance budget is more forgiving. Additionally, it has been shown that, even in fairly weak piezoelectric materials, piezoelectric coupling to bulk phonons in the substrate may be a significant loss channel that adversely effects delicate quantum circuits [179]. Thus, we initially opted

Design Parameters		Phenomenological Parameters	
Substrate	ST-X quartz	r_s	$-0.01i$
d	1 mm	Propagation loss	500 Np/m
x	250 μm	Derived Parameters	
z	4 μm	C_1	179 fF
N_p	25	C_J	9.8 fF
λ_{IDT}	600 nm	C_p	12.8 fF
w	70 μm	C_{SAW}	129.3 fF
λ_{mirror}	604 nm	f_0	5.203 GHz
N_{grat}	947	C_{eff}	75.1 fF
L_{free}	$\approx 3 \mu\text{m}$	$g_m/(2\pi)$	20.5 MHz

Table 6.2: Summary of the target design parameters, phenomenological parameters, and derived parameters for the device on ST-X quartz experimentally tested and described in § 6.4.1.

to work with ST-X quartz as the piezoelectric substrate, a fairly weak piezoelectric material, where acoustic losses would hopefully be mitigated.

The targeted design parameters, as well as the circuit/Hamiltonian values derived using the modeling method from § 6.2 are listed in Table 6.2. Note that the free propagation length L_{free} is only approximately defined: in order to maximize the number of Bragg gratings/reflectivity of the mirrors, the resonator was lithographically defined in three separate writing steps (one for the IDT and one for each mirror) with stage translations between each step. Because stage translations of the scanning electron microscope are only accurate to several microns, we can only give an approximate value for this number.

In this experiment, we will still use the electromagnetic 3D cavity to manipulate/measure the state of the coupled cavity/qubit system. One parameter left unconstrained by our model was the coupling g between the transmon-like mode of the qubit and the 3D cavity mode. We would like g to be large enough such that we are still in the strong coupling regime between the 3D cavity and the transmon mode, so that we can control and read out the qubit state using the same methods as previous chapters. To verify coupling between the 3D cavity

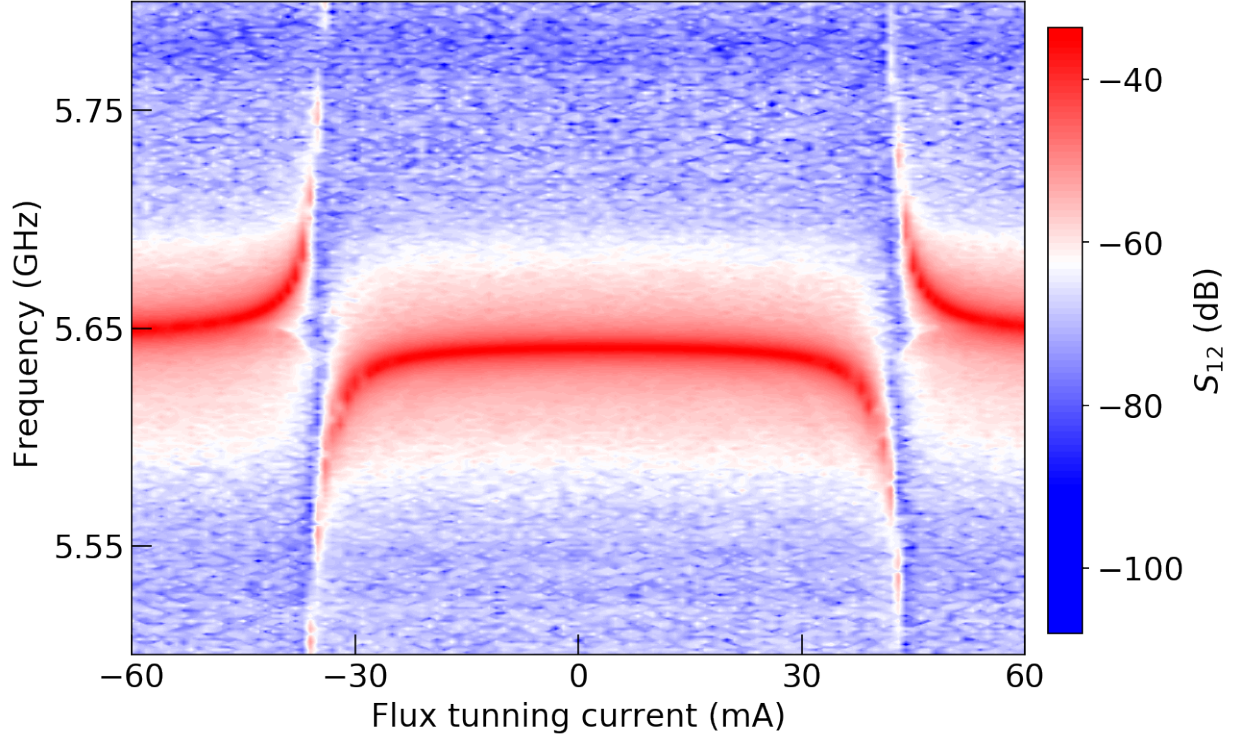


Figure 6.8: **Verifying coupling of the transmon-like mode to the electromagnetic cavity:** Transmission S_{12} through the 3D electromagnetic cavity with a qubit + ST-X quartz device loaded in, as a function of both frequency and flux tuning current.

and the transmon mode, we measure the transmission S_{12} through the cavity, as a function of applied flux bias. We see characteristic avoided crossings in the cavity spectrum as the qubit is flux tuned through the cavity resonance, roughly symmetric about zero flux bias, indicative of strong coupling between the 3D cavity and the flux-tunable transmon qubit mode. From this data, we may extract the cavity-qubit coupling rate $g/(2\pi) \approx 90$ MHz, which is of order the cavity-qubit coupling of a “standard” 3D transmon qubit.

Satisfied with the existence of a strongly coupled transmon-like mode in the cavity, we may now perform two-tone spectroscopy to measure ω_{01} and ω_{12} of the transmon-like mode, from which we may experimentally verify $\alpha = E_C = \omega_{01} - \omega_{12}$. Data measuring these values are plotted in Fig. 6.9(a-b). This data is taken at flux bias current $I = 0$: we record

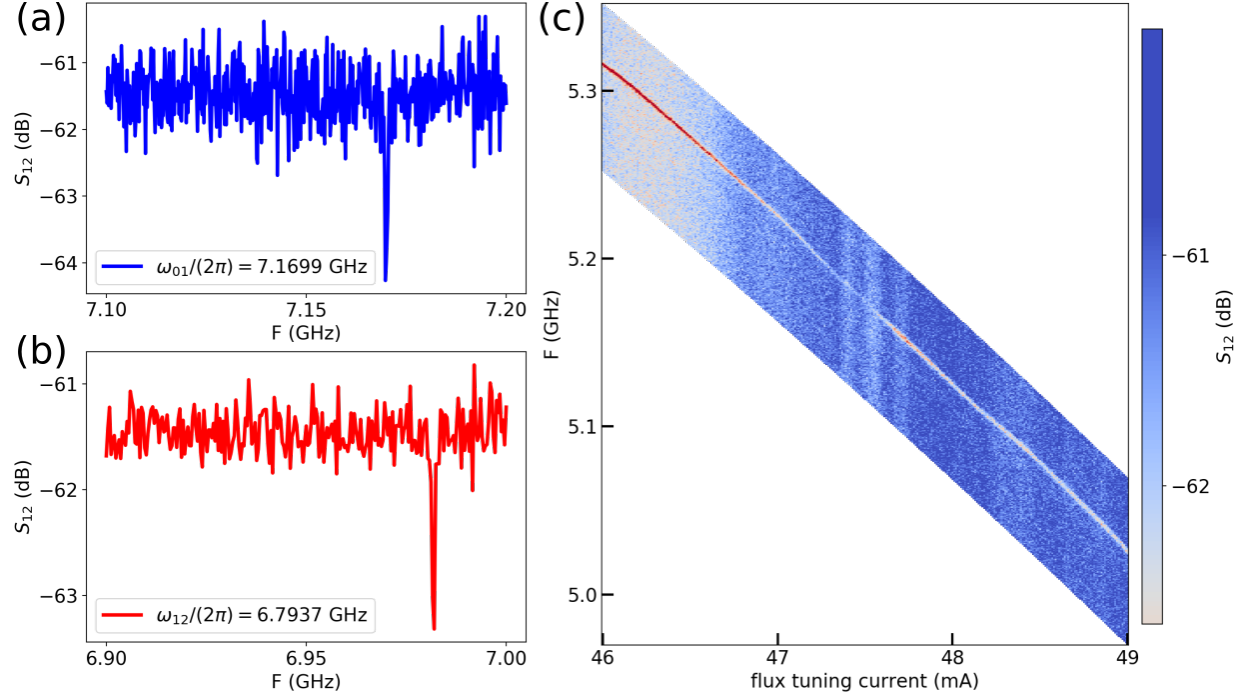


Figure 6.9: **Spectroscopy of the ST-X Quartz quantum acoustics device:** Two-tone spectroscopy measuring the (a) $|0\rangle \rightarrow |1\rangle$ transition and the (b) two-photon $|0\rangle \rightarrow |2\rangle$ transition at $I = 0$ flux bias current. (c) Two-tone spectroscopy, as a function of flux tuning current. As the frequency of the qubit changes, we raster the center frequency of the sweep to keep the $|0\rangle \rightarrow |1\rangle$ transition in the data range. We observe no obvious avoided crossings in the expected frequency range of the SAW device.

$\omega_{01}/(2\pi) = 7.170$ GHz and $\alpha/(2\pi) = 376$ MHz, corresponding to $C_{eff} = e^2/2\alpha = 51.5$ fF, well below the value predicted from the model.

Furthermore, we may perform two-tone spectroscopy as a function of flux tuning current to extract ω_{01} as a function of flux. We expect that, as we tune ω_{01} through the SAW resonant frequency, we should see an avoided crossing in the data, indicating coupling to the SAW device. In Fig. 6.9(c), we perform such a measurement: at each value of current, we raster the center frequency of the two-tone sweep to keep the $|0\rangle \rightarrow |1\rangle$ transition frequency within the range of the sweep, while not vastly oversampling. Sweeping from $I = 46$ mA to $I = 49$ mA, which corresponds to $\omega_{01}/(2\pi)$ ranging from ≈ 5.3 GHz to ≈ 5 GHz, we see no

evidence of an avoided crossing, which we would naively expect to be easily visible from the model outlined above.

Evidently, some portion of the model is deficient in describing the actual experiment in the fridge. In what follows, we explore some potential possibilities and suggest ways in which they may be accounted for.

6.4.1.1 Uncontrolled substrate spacing

That the capacitance of the transmon mode C_{eff} is substantially different from theory predictions is a strong indicator that spacing between the substrates in the flip-chip stack may be ill-controlled. Since all electrodes on a wafer are defined with high precision in lithography, and the pads of the two devices may be reliably aligned relative to one another, the only major adjustable parameter in determining the geometric capacitance between two electrodes is the height z between the substrates.⁵ To investigate this, we may perform finite element simulations at several heights z to obtain the values of capacitive network, as a function of z . Using these values, and all other values outlined in Table 6.2, we calculate the anticipated C_{eff} and g_m as a function of substrate spacing z , which we plot in Fig. 6.10(a-b). From these calculations, we expect the experimentally obtained value $C_{eff} = 51.5$ fF to occur if the substrates are $z \approx 10$ μm apart.

To investigate if this is in fact the case, we image the edge of the sample in a variable pressure SEM with tilt control. As is evident in Fig. 6.10(c), too much gluing resist was applied to this device, and the distance z between the substrates is limited by the glue resist and not the hard baked bump-bonds (visible as lines on the bottom substrate.) While it is hard to estimate the distance z from the angle of this image, it is certain that the spacing

⁵The substrates may also be tilted with respect to each other, breaking assumption of a symmetric capacitance network. We plan to explore the effect of substrate tilt in the future.

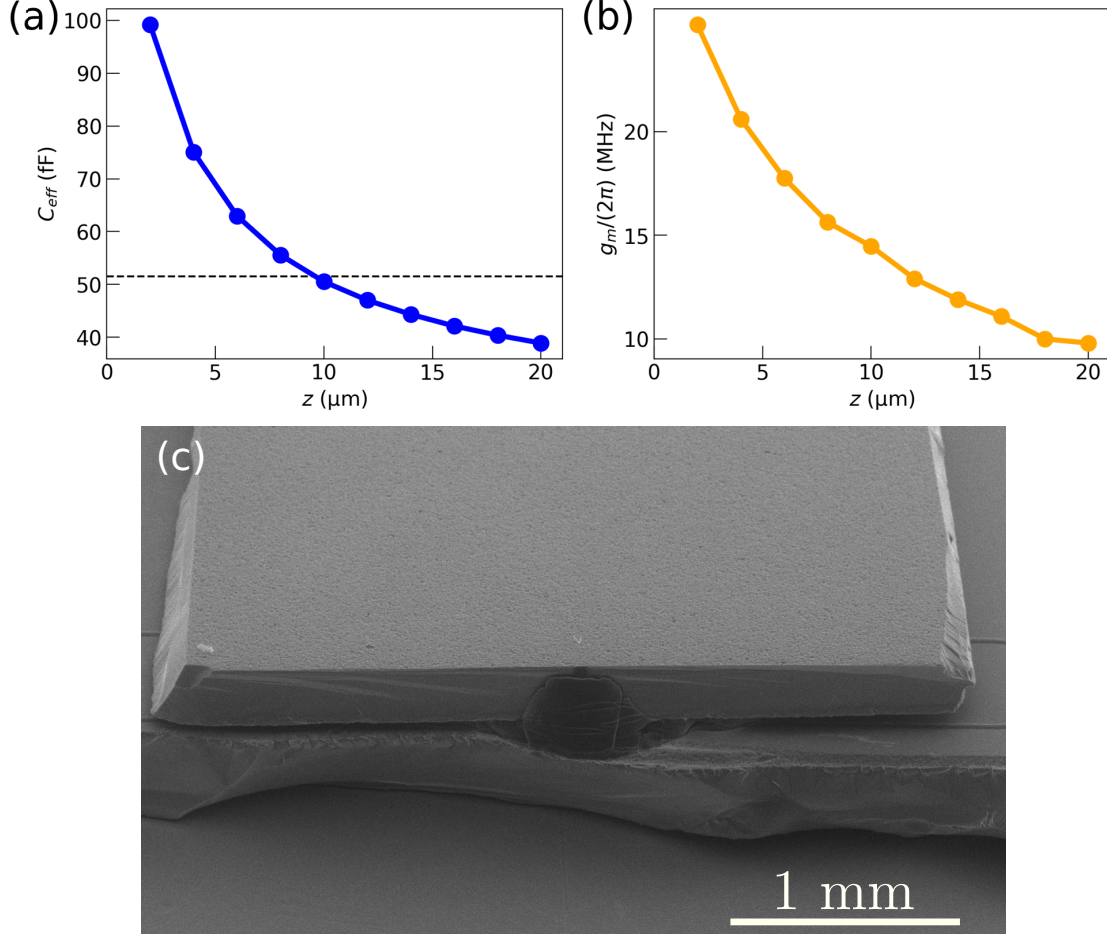


Figure 6.10: **Dependence on the substrate spacing z :** (a) Transmon mode capacitance C_{eff} and (b) acoustic coupling g_m as a function of height, using the all other parameters from the model outlined in Table 6.2. Empirically determined mode capacitance marked with a dashed black line. (c) SEM picture of the flip-chip stack from the edge, indicating that the spacing of this chip-stack limited by the glue rather than the hard baked spacers.

is $> 4 \mu\text{m}$. We have taken steps in the fabrication process to more accurately control the substrate spacing, however in the future it may be necessary to move away from soft baked resist applied by hand and to a different adhesive, or perhaps more precise substrate bonding technique such as indium bump bonding.

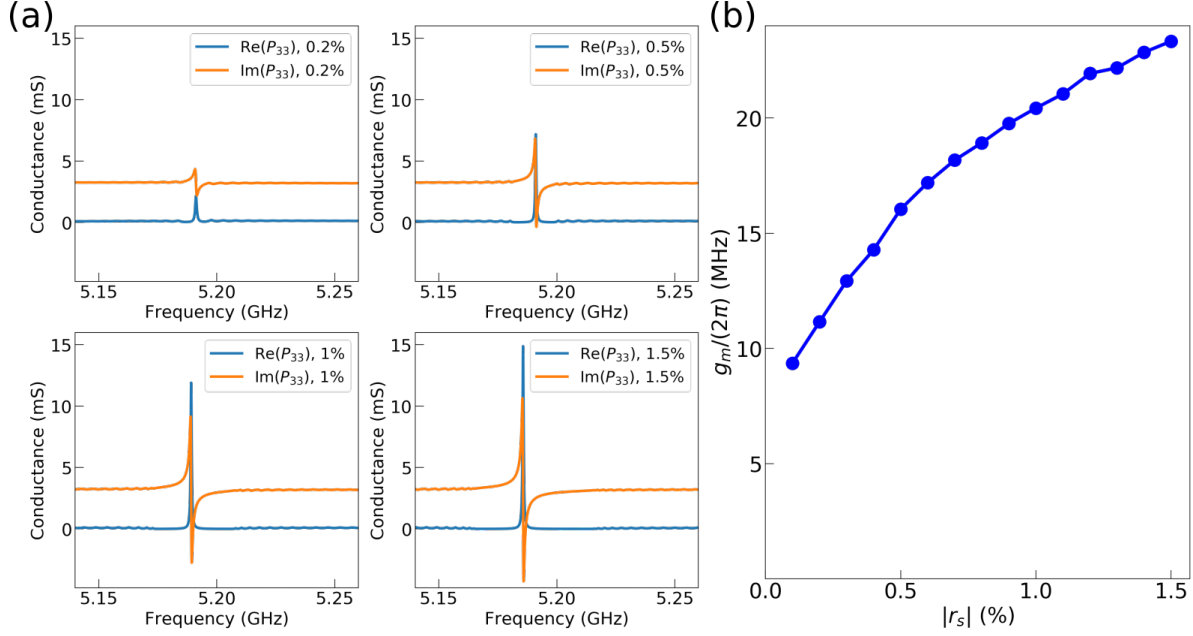


Figure 6.11: **Dependence of the SAW response on the grating reflectivity:** (a) Real (blue) and imaginary (orange) parts of the acoustic admittance $Y = P_{33}$, calculated using various values of the per-grating reflectivity. As the reflectivity decreases, the electrical response of the resonator becomes weaker. (b) g_m calculated at various values of $|r_s|$ using all other values defined in Table 6.2.

6.4.1.2 Mirror reflectivity

While the increased substrate spacing z will result in decreased acoustic coupling g_m , from Fig. 6.10(b) we see that the predicted coupling at $z = 10 \mu\text{m}$ is still of order $\approx 13 \text{ MHz}$, i.e. large enough to be resolved in our data. This is indicative that some assumptions we made in our model of the SAW resonator were faulty. The most obvious assumptions we made were in the reflectivity per Bragg grating r_s and the propagation loss. Propagation loss in SAW materials is ill-characterized at low temperature [167], however some studies [169] have recorded per-grating reflectivity as low as $|r_s| \sim 0.2\%$ on ST-X quartz at cryogenic temperatures. Given the large disparity between this and our guess of $|r_s| = 1\%$, we should ask what effect the reflectivity has on the acoustic coupling.

In Fig. 6.11(a), we plot both the real and imaginary parts of the acoustic admittance as

a function of frequency in the vicinity of the SAW resonance for several values of $|r_s|$. We observe that at low values of $|r_s|$, the acoustic admittance associated with the resonance is small, and the magnitude of the admittance grows as $|r_s|$ grows. From our analysis in Chapter 5, we know that as $|r_s|$ decreases the effective length L_{eff} of the SAW Fabry-Pérot cavity increases, and the IDTs will transduce SAWs/applied voltages over a smaller portion of the cavity. We also recall that the width of the mirror stop band Δf decreases as $|r_s|$ decreases, which can be observed Fig. 6.11(a) in as a decrease frequency range over which fine-scale features of $\text{Re}(P_{33})$ are suppressed.

Unsurprisingly, when the acoustic admittance of the resonator is suppressed, the coupling rate g_m to the qubit mode is also suppressed. In Fig. 6.11(b), we plot g_m calculated at several values of $|r_s|$. We see that below $|r_s| \sim 1\%$, the SAW-qubit coupling g_m begins to drop precipitously.⁶ Note that these calculations are done using all other values of the model recorded in Table 6.2, i.e. for $z = 4 \mu\text{m}$, and loss of coupling from low reflectivity will compound on the loss of coupling caused by unconstrained substrate spacing.

6.4.1.3 IDT position between the mirrors

Low per-grating reflectivity r_s could conceivably be accounted for if we could fabricate a very large array of Bragg gratings and simultaneously engineer the SAW resonance condition to be near the middle of the mirror stop band to optimize coupling. However, these goals turn out to be contradictory. Recall that the SAW resonator is written in three electron beam writing steps, with translations in between each, rendering the free propagation distance L_{free} only approximate. However, the acoustic coupling g_m turns out to be *exquisitely* dependent upon

⁶Since the gratings are finite length, at some low $|r_s|$ the quality factor of the modeled resonator, previously limited by propagation loss, starts to drop off. This will presumably also decrease coupling, though a full analysis along with the dependence on the propagation loss is needed.

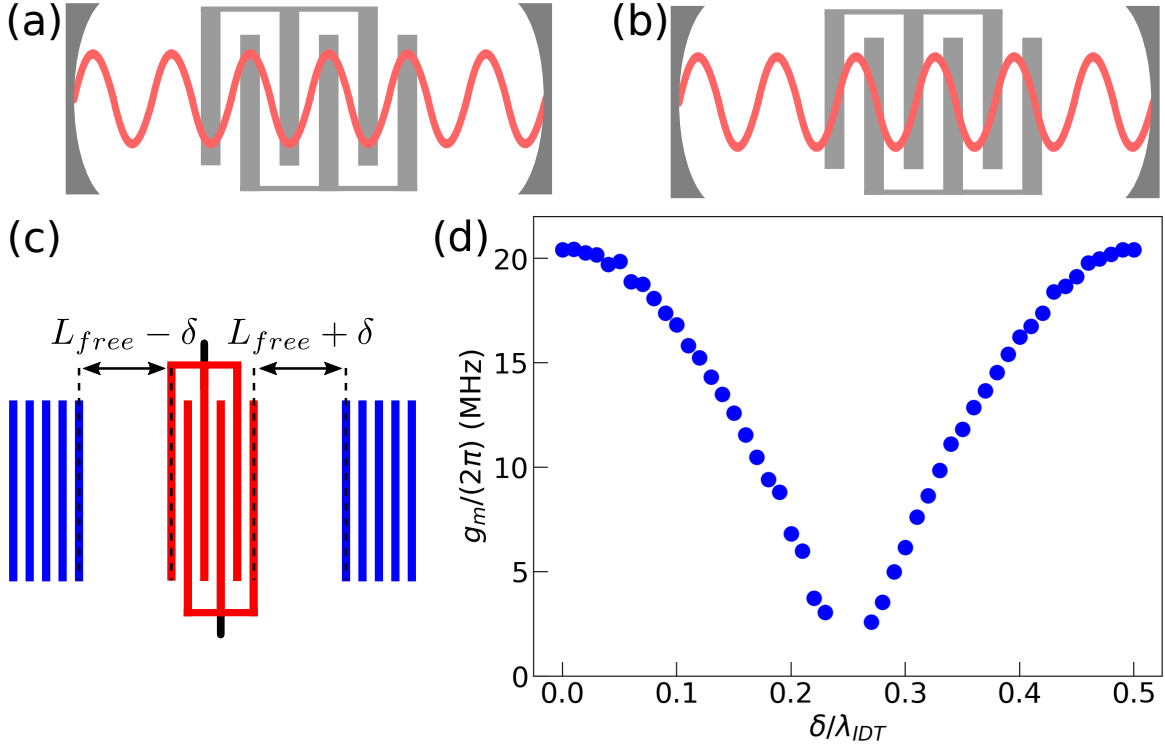


Figure 6.12: **Dependence of g_m on the IDT position:** Acoustic coupling is maximized when the IDT fingers are in phase with the antinodes (a) of the standing wave pattern of the Fabry-Pérot resonator, and is minimized when the IDT fingers are in phase with the nodes (b). (c) To investigate this effect, we model the resonator and shift the position of the IDT inside of the fixed-length resonator by δ . (d) Acoustic coupling g_m vs δ . As expected, coupling is maximized when δ is a half-integer value of λ_{IDT} , and minimized at odd-quarter integer values. Note that near $\delta/\lambda_{IDT} = 0.25$, the response becomes too small to accurately model. Here, $\delta = 0$ corresponds to the model defined by the values in Table 6.2.

the free propagation distance.

To see this, recall that the resonance condition in a Fabry-Pérot resonator is set when the total resonator length L_{eff} equals an integer number of half wavelengths, i.e. it is set by a standing wave condition. This standing wave will have nodes/antinodes at fixed points in space, which the IDTs may be in-phase/out-of-phase with, as depicted in Fig. 6.12(a-b). If the IDT is in-phase with the antinodes, as in Fig. 6.12(a), the voltage induced across the IDT electrodes will be maximized, and transduction will be maximized. Conversely, if the IDT is in phase with the nodes, in Fig. 6.12(b), a confined SAW will not cause a voltage

will build up across the IDT electrodes, and transduction will effectively drop to zero. This is outlined in Fig. 6.12(d), where we calculate the acoustic coupling to the transmon mode g_m as a function of IDT offset δ from the center of the resonator. Around $\delta/\lambda_{IDT} = 0.25$ (where the IDT fingers are in phase with the standing wave nodes) the acoustic coupling drops precipitously.

For small wavelength devices, such as the device cooled down with $\lambda_{IDT} = 600$ nm, this presents a difficult fabrication problem: the IDT and mirrors (each write of scale $\sim 50 - 500$ μm , with the whole structure close to 1 mm) must be aligned relative to each other with an accuracy $\ll \lambda_{IDT}/4$ in order to have predictable acoustic coupling. While this is perhaps possible, accounting for imperfection in the electron beam lithography calibration and thermal contractions when cooling down to 10 mK make it a daunting engineering task. It would be preferable to write the entire structure (IDT + mirrors) in a *single* write, however we are only capable of writing finite size structures while maintaining accuracy at the length scale required to define the IDT/gratings. Pragmatically, this limits the number of gratings N_g each mirror can have, which will in turn limit the reflectivity of the mirrors and the quality factor of the SAW resonator, especially for structures fabricated on weak piezoelectric substrates such as ST-X quartz where $|r_s|$ may be low.

6.4.1.4 Coherence times

We have also recorded coherence times T_1 and T_2 at multiple values of flux tuning current, ranging from 47-48 mA. Over the range of frequencies measured in this device, we recorded an average $T_1 = 1.77$ μs , and average $T_2 = 388$ ns. This device (both SAW resonator and Josephson junction device) were fabricated completely in-house, and the observed value of T_1 is fairly consistent with other “standard” 3D transmon that were fabricated at MSU.

This is consistent with the picture that, in the case of quantum acoustics devices employing ST-X quartz, the energy relaxation rate is currently limited by our fabrication capabilities. The low T_2 is not terribly surprising: we are in a regime where the qubit frequency is strongly dependent on the applied flux, meaning the qubit is highly susceptible to flux noise. Additionally, and perhaps more importantly, the low C_{eff} means that, in this regime $E_J/E_C \sim 27$ i.e. the qubit is in the “offset charge sensitive” transmon regime, where the transition frequency ω_{01} experiences a nontrivial amount charge dispersion and fluctuations in the local electrostatic environment may drive dephasing. Improving flux noise filtering and increasing the capacitance into the full-fledged transmon regime ($E_J/E_C \geq 50$) would likely improve T_2 .

6.4.2 Devices on LiNbO₃

Given the above considerations, the natural course of action is to move to a material with stronger piezoelectricity, such as LiNbO₃. This choice was not originally pursued due to concerns that the piezoelectricity would be *too* strong, such that microwaves in 3D cavity would be transduced into spurious bulk modes in the LiNbO₃ substrate, opening up a loss channel that would render the experiment inoperable. These concerns were facilitated by early tests, where the decay rate κ of a 3D cavity increased significantly when a LiNbO₃ substrate was placed inside. Fig. 6.13 shows the transmission S_{12} through a 3D cavity with/without a LiNbO₃/silicon flip-chip stack: in the presence of LiNbO₃, κ increases significantly, and $S_{12,max}$ decreases significantly. However, these concerns are alleviated at low temperatures, where the loss rate κ significantly *decreases*, even compared to the no device configuration. It is unclear why this is case, though there may be some ensemble of defects/contamination on LiNbO₃ that freeze out at low temperatures. A summary of these results is found in

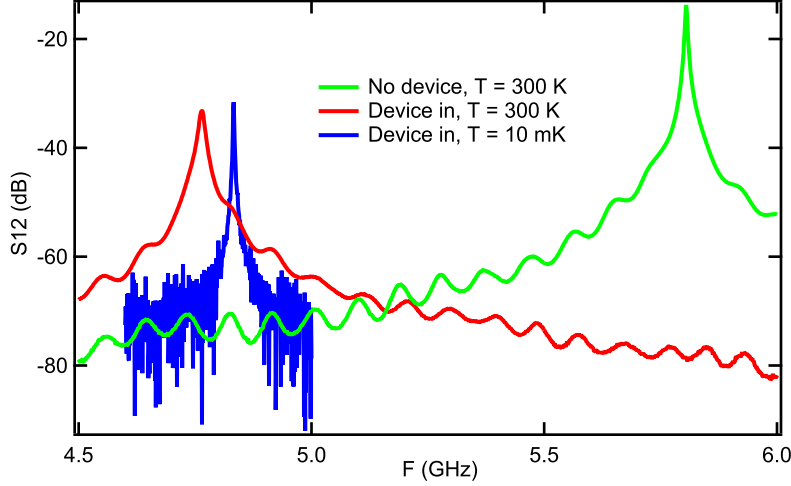


Figure 6.13: **3D cavity transmission in the presence of LiNbO₃**: Transmission of a 3D cavity at room temperature with no device (green), at room temperature with a LiNbO₃/silicon flip-chip stack (red), and at ~ 10 mK with a LiNbO₃/silicon flip-chip stack (blue). Note that the absolute magnitude of the low temperature also encapsulates the experimental attenuation/amplification, and should therefore not be directly compared to the magnitude of the other data.

Configuration	f_0	κ
Room temperature, no device	5.8049 GHz	4.8 MHz
Room temperature, device in	4.7644 GHz	13.3 MHz
10 mK, device in	4.8327 GHz	2.2 MHz

Table 6.3: Transmission characteristics of a 3D cavity with/without LiNbO₃, extracted from the data plotted in Fig. 6.13

Table 6.3.

We have begun to cool down quantum acoustics experiments that use LiNbO₃ in recent months. While we have not yet cooled down a device in a parameter regime in which we would expect acoustic coupling⁷, we have taken data, qualitatively similar to the data presented in Figures 6.8 and 6.9, indicating once again the presence of a well defined transmon qubit mode in the experiment. In this device, we measured $\alpha = E_C = 236$ MHz, corresponding to

⁷The most recent experiment had a transmon qubit with an unexpectedly low E_J , such that $\omega_{01} < \omega_m$ even at zero flux bias.

$C_{eff} = 82$ fF, more in line with what we expected from simulations.

We have also measured the coherence times of the transmon-like modes for devices fabricated with LiNbO₃. Preliminary measurements of T_1 on devices employing LiNbO₃ have produced $T_1 \approx 140$ ns, roughly an order magnitude lower than the case of ST-X quartz, indicating that there is perhaps significant material loss in LiNbO₃ at the single photon level. We stress that this data is preliminary, and that coherence times of superconducting qubits are known to vary widely for a number of (sometimes ill-controlled) reasons: more definitive comments on the loss of LiNbO₃ would require repeating the experiment several times. Preliminary measurements of T_2 for devices on LiNbO₃ have shown $T_2 \approx 240$ ns, i.e. close to the relaxation limited value $T_2 = 2T_1$. This is somewhat unsurprising, since for this device $E_J/E_C \approx 48$, squarely in the transmon regime where charge noise is exponentially suppressed. This indicates that there are likely no additional sources of pure dephasing in LiNbO₃, though we again stress that these values are preliminary and more thorough measurements are required.

6.5 Future work

While we are still working to demonstrate coupling between SAW resonators and 3D transmon qubits, there are several parallel paths that are worth exploring, several future improvements that could be imagined, and several experiments that could be performed once coupling has been established. We outline a few of them in this section.

As we have seen, a pressing issue in correlating the model to real experiments remains the uncontrolled phenomenological parameters in the model. For a given quantum acoustics device, our best bet approach would be to fit any recorded data to the model and extract

these parameters, however for design purposes it would be advantageous to know them before hand. Systematic studies of SAW structures in planar geometries, such as in Ref. [169], could help determine the per-grating reflectivity and propagation loss at low temperatures. We have already begun taking steps to perform these studies, and plan to do so in the future.

While not surprising, evidence of suppressed T_1 for devices employing LiNbO₃ indicates that devices on strongly piezoelectric substrates may have finite utility. One strategy to resolve this issue could be to greatly reduce the amount of LiNbO₃ in the system: recently, techniques have been developed to create single crystalline thin film LiNbO₃, with thicknesses of order several μm bonded to the surface of a silicon wafer. These films have been shown to host Rayleigh waves [180], and could greatly reduce the available density of states for the qubit to piezoelectrically decay into, allowing us to maintain both strong coupling and long coherence times.

An interesting feature of classical SAW devices is the ability to engineer low dissipation structures with responses that have sharp features as a function of frequency, such as band pass filters [160]. These sharp frequency response structures may be a valuable resource for quantum devices. For example, one could imagine creating a SAW based Purcell filter [181], i.e. a filter that suppresses the density of states for a qubit to decay into. Additionally, the sharp frequency response of SAW devices could be used to engineer the dissipation environment of the qubit, preferentially increasing dissipation at some frequencies and suppressing it at others. This type of “bath engineering” has been used to autonomously prepare arbitrary states of a qubit system [182] and to create dissipation hierarchies that simulate the response of non-hermitian Hamiltonians [183].

Chapter 7

Acoustoelectric transport in graphene

Long before surface acoustic waves were incorporated into quantum circuits, SAWs were used extensively in condensed matter physics to study systems on or near the surface of piezoelectric substrates. These techniques were particularly successful in the semiconductor community, where the AlGaAs/GaAs heterostructures that host two-dimensional electron systems (2DESs) are naturally fabricated close to the surface of a piezoelectric GaAs substrate. Many interesting properties of the integer and fractional quantum Hall states of matter have been elucidated by SAW techniques, most notably experimental evidence for the formation of a composite fermion surface at $\nu = 1/2$ [184, 185].

Much of the motivation for the main focus of this thesis, superconducting qubits and quantum acoustics, was generated by experiments done earlier in my Ph.D. studying the interaction of SAWs with another 2DES: exfoliated graphene. This chapter catalogs those experimental efforts, which were approached with the quantum Hall effect in mind. Thus while none of the data presented is taken in high magnetic field, much of the motivation will be given through the lens of previous experiments using SAWs to study the quantum Hall state, and emphasis will be given to optimizing the experimental setup for high magnetic field measurements. An introduction to the quantum Hall effect will not be given: the interested reader may find an excellent introduction in Ref. [186].

While this chapter stands out thematically from the rest of the thesis, I feel that it is

important to include it nonetheless. In section 7.4, I will briefly discuss how some of the ideas developed in this chapter may be integrated with the experimental ideas pertaining to superconducting qubits and quantum acoustics presented in the rest of this thesis. The results in this chapter were also reported in Ref. [187].

7.1 Graphene: a brief introduction

Graphene is an allotrope of carbon consisting of a single layer of carbon atoms arranged in a two-dimensional hexagonal lattice. One may think of graphene as the single atomic layer limit of graphite, or conversely think of graphite as many layers of graphene stacked on top of each other. The electronic properties of graphene were first predicted in 1947 [188], however it would not be until 2004, over 50 years later, that monolayer graphene was isolated and its electron transport properties were experimentally confirmed [189]. This long gap elicits the image of a technologically complex method of isolating graphene, however reality often turns out to be stranger: graphene was first isolated by mechanically exfoliating graphite flakes and depositing them onto oxidized silicon substrates using Scotch tape!¹ Indeed, to this day the highest mobility graphene devices, while vastly more complicated than early devices, are still made by mechanical exfoliation using tape [190].

The ability to mechanically exfoliate and measure electrical transport through an intrinsically 2D material has given experimentalists a highly flexible toolkit to engineer a wide variety of experiments studying the 2D electron system formed by the charge carriers in graphene. The choice of oxidized silicon as a substrate to suspend the graphene flakes was prudent: the

¹A quick internet search pegs the price of a dispenser of Scotch tape at $\approx \$2.30$. Given that the 2010 Nobel Prize was awarded for the isolation of graphene, by my calculations, Scotch tape is by cost roughly 4×10^8 times more efficient at producing Nobel Prizes than the LHC.

oxide layer forms an insulating barrier between the graphene and the highly doped silicon bulk, allowing for experimentalists to modulate the Fermi energy of the graphene *in situ* simply by applying an electrostatic voltage to the conducting host substrate [189]. Quickly after graphene was isolated, the remarkably high mobility of exfoliated devices led to the observation of the integer quantum Hall effect in graphene [191, 192]. Further advances in fabrication techniques, such as encapsulating graphene in mechanically-exfoliated insulating hexagonal boron nitride [193] have pushed up the mobility of layered 2D devices even further, allowing for the observation of both odd [194, 195] and even [196, 190] denominator fractional quantum Hall states. Placing multiple exfoliated flakes with crystal axes misaligned, or “twisted” relative to each other, has produced a bounty of experiments including the observation of Hofstadter’s butterfly [197] and (possibly unconventional) superconductivity in bilayer graphene [198].

In addition to its fundamentally 2D nature, the band structure of graphene leads to a wide variety of interesting physics. Graphene is a Dirac semimetal, with the valence and conduction bands touching at a single point at the Fermi level. Near the Fermi energy, the bands are linear as a function of momentum ($E = \hbar k v_f$ where v_f is the Fermi velocity), rather than the typical approximation of quadratic bands at low energies ($E = \hbar^2 k^2 / 2m^*$, where m^* is the effective mass of the charge carriers.) This linear dispersion relation implies that charge carriers in graphene are effectively *massless* quasiparticles with an effective “speed of light” equal to the v_f . This massless nature of the low-energy excitations means that the charge carriers in graphene are better described by the *Dirac* equation, rather than the Schrödinger equation, and has led to a number of observations of “relativistic” phenomena. Prominent among these is the observation of Klein tunneling, where a massless fermion paradoxically tunnels through an infinite potential well [199] and the quadratic spacing of

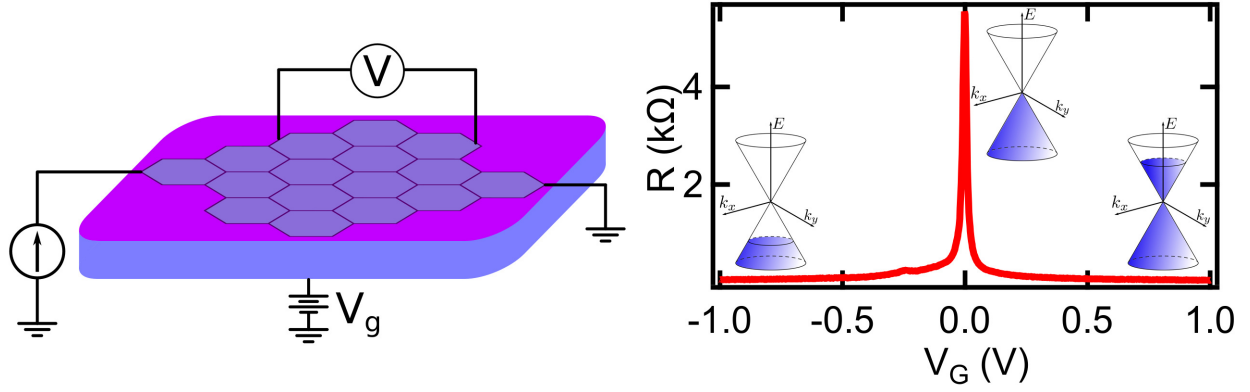


Figure 7.1: **DC transport measurements of graphene:** Left: Schematic of a generalized DC transport measurement of graphene. Graphene (turquoise hexagons) is isolated and stabilized on the surface of an oxidized silicon substrate. The bulk of the substrate (blue) is heavily doped, such that the silicon remains conducting at low temperatures, while the graphene sample is electrically isolated from the conducting bulk by a several hundred nanometer thick oxide layer (purple.) The conducting graphene sheet and the bulk silicon form a capacitor, and applying a DC gate voltage V_G to the silicon varies the charge in the graphene, which modifies the position of the Fermi energy. Right: resistance of the graphene sheet as a function of gate voltage. At low and high gate voltages, the Fermi energy sits at positions in the band with a high density of states (left and right insets) resulting in a low sample resistance. When the Fermi energy is tuned close to the Dirac point, where the conduction and valence bands meet and the density of states goes to zero (middle inset), the resistance increases dramatically. Note that the type of charge carrier changes from holes to electron as the Fermi energy is tuned from the conduction band to the valence band. Data taken by Liangji Zhang, and reprinted with permission.

Landau levels for massless particles, which extraordinarily has led to the observation of the integer quantum Hall effect at room temperature [200].

7.2 Surface acoustic waves and two-dimensional electron systems

As elaborated in Chapter 5, SAWs propagating on piezoelectric substrates carry with them a co-propagating electric field with the same frequency and wave vector as the SAW, and this electric field extends evanescently above the surface of the bulk with decay length λ , the

wavelength of the SAW. When the charge carriers of a nearby conductor are in the vicinity of this co-propagating electric field, the SAW and the charge carriers interact, modulating the SAW propagation and moving the charges. The interaction of SAWs and charge carriers in two-dimensional electron systems close to the surface of propagation has been employed in many fundamental studies and proposed devices. Here, we sample two classes of studies commonly executed: velocity shift/attenuation studies, where the conducting electrons screen the SAW and modulate its propagation, and acoustoelectric studies, where a propagating SAW electric field imparts momentum on the charge carriers in the 2DES.

7.2.1 SAW attenuation and velocity shift

If a SAW propagates in the vicinity of a conductor, the co-propagating electric field will induce currents in the conductor, leading to Ohmic dissipation. This will attenuate the SAW as it passes by the conductor (and, as a consequence of the Kramers-Kronig relations, shift its propagation velocity.)² The velocity² and attenuation per unit length of a SAW propagating on a piezoelectric substrate with conductivity σ may be written as [201, 202]

$$v(\omega) = v \left(1 + \frac{K_{eff}^2/2}{1 + (\omega/\omega_c)^2} \right) \quad (7.1a)$$

$$\Gamma = k \frac{K_{eff}^2}{2} \frac{(\omega/\omega_c)}{1 + (\omega/\omega_c)^2} \quad (7.1b)$$

Here k is the SAW wave vector, v_0 is the SAW velocity in the limit where $\sigma \rightarrow \infty$, K_{eff}^2 is the effective piezoelectric coupling constant that depends on material parameters

²Unfortunately, the convention within the 2DEG community is to reference the velocity with respect to the perfectly conducting surface, as opposed to our discussion in Chapter 5 where we referenced it to an insulating surface. I have decided to keep in line with the quantum Hall literature here, so in this chapter, velocity shifts caused by a conducting layer at the surface will be positive rather than negative.

and geometry³, and $\omega_c = \sigma/(\epsilon_1 + \epsilon_2)$, where $\epsilon_{1(2)}$ is the dielectric constant above (below) the surface. ω_c may be thought of as the inverse of a time constant that quantifies how quickly charges can equilibrate after being perturbed by an external electric field. In the thin film limit, where the conducting film thickness is much smaller than the SAW wavelength $d \ll k^{-1} = \lambda$, the effective conductivity of the conducting layer (as sampled by the SAW electric field) reduces by a factor of $dk \ll 1$ [202]. We may then write

$$\omega_c = \frac{\sigma dk}{(\epsilon_1 + \epsilon_2)} = \frac{\sigma_{\square} k}{(\epsilon_1 + \epsilon_2)}$$

where $\sigma_{\square}$ is the sheet conductivity.

In practice, the SAW frequency is more or less fixed by choice of transducer geometry, so it is much more convenient to write the velocity and attenuation in terms of the conductivity of the 2DES, which we can often vary. If we do so, we arrive at

$$\Delta v/v = \frac{K_{eff}^2}{2} \frac{1}{1 + (\sigma/\sigma_m)^2} \quad (7.2a)$$

$$\Gamma/k = \frac{K_{eff}^2}{2} \frac{(\sigma/\sigma_m)}{1 + (\sigma/\sigma_m)^2} \quad (7.2b)$$

where $\sigma_m = \omega/k(\epsilon_1 + \epsilon_2) = v_0(\epsilon_1 + \epsilon_2)$ is a material dependent characteristic conductivity. Here, we have taken the experimentally expedient step of writing down the velocity as the relative shift referenced to the case of a perfect conductor. Importantly we see that the velocity shift and attenuation are a function of both the SAW frequency and wave-vector, implying that these measures probe the frequency and wave vector dependent conductivity

³For a conducting layer directly on the surface of the piezoelectric, $K_{eff}^2 = K^2$, the piezoelectric coupling constant defined in Chapter 5. From Eqn. 7.1a, it may be verified that $v(\sigma = \infty) - v(\sigma = 0) = K_{eff}^2/2$.

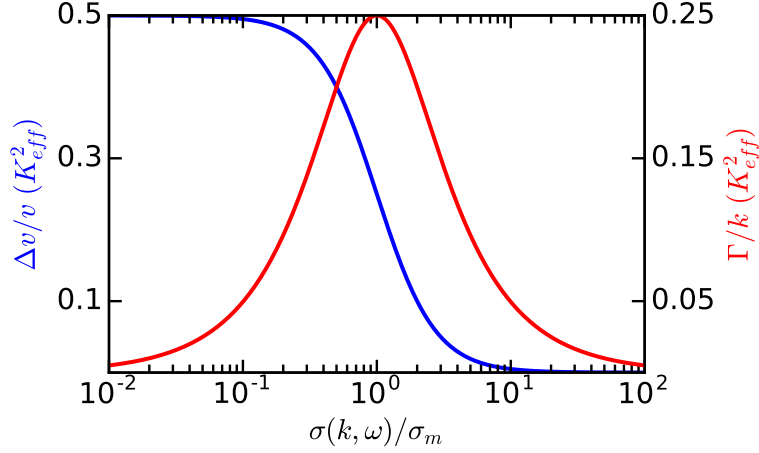


Figure 7.2: **SAW attenuation and velocity shift vs. 2DES conductivity:** Relative SAW velocity shift $\Delta v/v$ (blue, left axis) and attenuation per unit length normalized by the SAW wave vector Γ/k (red, right axis) as a function of normalized 2DES conductivity $\sigma(k, \omega)/\sigma_m$. Both are given in units of the effective piezoelectric coupling constant K_{eff}^2 .

$\sigma(k, \omega)$. The normalized values of Γ/k and $\Delta v/v$ are plotted in Fig. 7.2.

Velocity shift and attenuation techniques are most sensitive in probing large changes in the bulk 2DES conductivity around σ_m , which tends to be low (for a conducting film on the surface of {110} GaAs, $\sigma_m \approx 3 \times 10^{-7} (\Omega/\square)^{-1}$ [202]). In the quantum Hall regime, current is dominated by conducting edge channels, and transport characterization of the insulating bulk is difficult. Here, SAW techniques have proved particularly valuable, and have been employed extensively to study high mobility 2DESs in GaAs/AlGaAs heterostructures in the quantum Hall regime [203, 202, 204, 205, 206, 207]. In particular, the wave vector dependence on the conductivity proved crucial for providing evidence for the existence of a composite fermion metal near Landau level filling factor $\nu = 1/2$ [184, 185]. SAW techniques have also been used in GaAs/AlGaAs heterostructures to study Wigner crystallization at low carrier density [208, 209, 210], and the 2D metal-to-insulator transition at zero magnetic field [211].

7.2.2 Acoustoelectric effect

Conversely, the electric field co-propagating with the SAW may impart momentum on the charge carriers in the 2DEG, driving a macroscopic DC current [212, 204, 213, 214, 215]. We present a rough sketch of the physical picture that follows that of Ref. [213]. Consider a SAW propagating in the x -direction with a co-propagating plane-wave electric field

$$E_x(x, t) = E_0 e^{i(kx - \omega t)} \hat{x} \quad (7.3)$$

If the SAW power isn't too large, this electric field will generate a *local* oscillating Ohmic current density $j_\alpha(x, t)$ direction 2DES

$$j_\alpha(x, t) = \sigma_{\alpha x} E_x(x, t) \quad (7.4)$$

where $\alpha = (x, y)$ is the component of the acoustoelectric current in question and σ_{ij} is the (i^{th}, j^{th}) component of the conductivity tensor of the 2DES. We may equivalently describe this current as an oscillation of the charge density N_s about its equilibrium value $N_{s,0}$

$$N_s(x, t) = N_{s,0} + \Delta N_s e^{i(kx - \omega t)} \quad (7.5)$$

Once again, assuming the SAW power isn't too high, $\Delta N_s \ll N_{s,0}$ and we may expand the conductivity as a function of the carrier density

$$\sigma(x, t) = \sigma_0 + \frac{\partial \sigma}{\partial N_s} \Delta N_s e^{i(kx - \omega t)} \quad (7.6)$$

We may then use the 1D continuity equation $\partial(-eN_s)/\partial t + \partial j_x/\partial x = 0$ write down ΔN_s in

terms of the SAW electric field

$$\Delta N_s(x, t) = -\frac{k\sigma_{xx}}{e\omega} E_x(x, t) = -\frac{\sigma_{xx}}{v_0 e} E_x(x, t) \quad (7.7)$$

We may then use this value to calculate the modified conductivity with Eqn. 7.6, and plug that back into Ohm's law to find the time averaged current $\langle j_\alpha(x) \rangle$ to first order⁴ in ΔN_s

$$\begin{aligned} \langle j_\alpha(x) \rangle &= \langle \sigma_{\alpha x}(x, t) E_x(x, t) \rangle = \left\langle \left(\sigma_{\alpha x, 0} + \frac{\partial \sigma_{\alpha x}}{\partial N_s} \Delta N_s e^{i(kx - \omega t)} \right) E_x(x, t) \right\rangle \\ &= -\frac{\partial \sigma_{\alpha x}}{\partial N_s} \frac{1}{v_0 e} \langle \sigma_{xx, 0} E_x^2(x, t) \rangle \end{aligned} \quad (7.8)$$

Since the time average of $E_x(x, t) = 0$, we can drop the first term, however the second term is proportional to $E_x^2(x, t)$, whose time average is *not* zero: this is the term that causes macroscopic acoustoelectric currents. In fact, we recognize that $\langle \sigma_{xx, 0} E_x^2(x, t) \rangle$ is the average Ohmic power dissipated by the charge carriers in the 2DES, and may write down the attenuation of the SAW per unit length as a result of the charge carriers in the 2DES as

$$\langle \Gamma \rangle = -\frac{1}{I} \frac{dI}{dx} = \frac{1}{I} \langle \sigma_{xx, 0} E_x^2(x, t) \rangle \quad (7.9)$$

where we have assumed the intensity⁵ of the SAW may be written as $I(x) = I_0 e^{-\Gamma x}$. We may then combine Equations 7.8 and 7.9 to write down the time averaged acoustoelectric current as a function of the SAW attenuation

⁴The condition for which this first order approximation is valid, i.e. when we can drop the term $\propto \Delta N_s^2$, is $2\mu|E_x| \ll v_0$, where μ is the carrier mobility. This condition is valid in the data presented in this chapter, where the mobility of the graphene is low, and wavelength of the SAW is large (and thus $E \propto \nabla\phi$ is smaller.) However the same may not be true of high mobility samples (such as those exhibiting the FQHE) at GHz frequencies.

⁵Here, intensity has units power/length.

$$\langle j_\alpha(x) \rangle = -\frac{\partial \sigma_{\alpha x}}{\partial N_s} \frac{1}{v_0 e} \langle I\Gamma \rangle \quad (7.10)$$

In the absence of a magnetic field, the off-diagonal components of the conductivity tensor disappear, and only remaining component of the time averaged current can be written as

$$\langle j_x(x) \rangle = -\frac{\partial \sigma}{\partial N_s} \frac{1}{v_0 e} \langle I\Gamma \rangle = -\frac{\mu}{v_0} \langle I\Gamma \rangle \quad (7.11)$$

where $\mu = N_s e \sigma$ is the carrier mobility in the 2DES. Thus, we see that, in the absence of a magnetic field, the *ac electric field* associated with the SAW drives a *dc current* in the 2DES in the same direction as the SAW propagation. The 2DES may also be subject to an external electric field in the $\hat{x}$ direction: in this case, we can write the overall DC current in the sample as

$$j = \sigma E - \frac{\mu I\Gamma}{v_0} \quad (7.12)$$

Similar to $\Delta v/v$ and attenuation measurements, acoustoelectric measurements have been employed to study in the quantum Hall effect in GaAs quantum wells, facilitated by the piezoelectricity of the host substrate [212, 204, 213, 214, 216]. More recently, work has been done using the acoustoelectric effect to pump *single* electrons, which has opened up the possibility of using SAWs as a mechanism for coherent control in quantum dot based quantum computing schemes [217, 218]. Additionally, the observation of plateaus in the acoustoelectric current through a 1-D channel [219, 220] has raised possibility of using the acoustoelectric effect as a metrological current standard. While the quantization of current plateaus in GaAs quantum wells remain insufficient for a current standard [221] there is some

hope that a cleaner system, such as electrons on helium [222], may exhibit highly quantized current plateaus.

7.3 Acoustoelectric effect in exfoliated graphene

In anticipation of doing experiments in the quantum Hall regime, we would like to use exfoliated graphene devices, which in general have much higher carrier-mobilities than are achievable than in, say, chemical vapor deposition (CVD) grown graphene. However, exfoliated graphene devices tend to be small (on the order ~ 10 's of μm or smaller) which creates an inherent problem for $\Delta v/v$ and attenuation measurements: the signal from these measurements are inherently *integrated* over the SAW propagation path. We measure the attenuation of the propagating SAW per unit length, or we measure the accumulated phase shift of a velocity shifted SAW: either way, small sample sizes severely limit the signal size. While it may be possible to design experiments around this (see Section 7.4 for some discussion of this), for first-order experiments it is much easier to simply focus on acoustoelectric measurements, where signals are inherently intensive.

In recent years several theoretical [223, 224, 225, 226] and experimental [227, 228, 229, 230, 231, 232] efforts have been made to extend SAW techniques to study the electronic properties of graphene. However, even in acoustoelectric studies, experimental challenges arise when applying these methods to graphene that are absent from similar GaAs experiments. The most pressing is the need to incorporate a compatible gate electrode for charge carrier density control. Oxidized silicon, the most common substrate for high quality graphene devices, is not piezoelectric, and the most strongly piezoelectric materials are also highly insulating, making gating graphene devices on these substrates difficult.

In this chapter, we report the observation of an *in situ* gate-controlled acoustoelectric effect in an exfoliated graphene device fabricated on an oxidized silicon substrate. To achieve this, we use a flip-chip device geometry, qualitatively similar to the quantum acoustics geometry from Chapter 6, where SAWs propagate on a separate piezoelectric substrate which is flipped upside-down and mechanically clamped to a silicon substrate holding the graphene. In this geometry, the electric field associated with the SAW evanescently extends above the LiNbO₃ surface, coupling to charge carriers in the graphene. We observe a clear dependence of this acoustoelectric voltage as we tune the sign of the charge carriers in the graphene.

7.3.1 Cold-finger and sample stage

One of the dilution refrigerators in our lab (a BlueFors LD-400) comes equipped with a 14 T superconducting solenoid magnet. By virtue of the solenoid not being infinitely long, the generated magnetic field is spatially inhomogeneous along the magnet bore axis, with a maximum field ~ 41.5 cm below the MXC. In anticipation of quantum Hall experiments, which necessitate mK temperatures and large magnetic fields, we built a “cold-finger” extension to the MXC plate from which we could mount exchangeable sample stages that sit at the height of the maximum magnetic field. A picture of this cold-finger is provided in Fig. 7.3(a).

The cold-finger is designed to minimize heating from the changing magnetic field, caused by either induced eddy currents or nuclear spin alignment [89], and to minimize the thermal resistance between the sample stage and the MXC. To efficiently thermally connect the sample stage and the mixing chamber plate, we used two 1/4”-diameter high-purity silver rods (Fig 7.3, (4)), which were annealed at 600° for 90 hours to remove grain boundaries

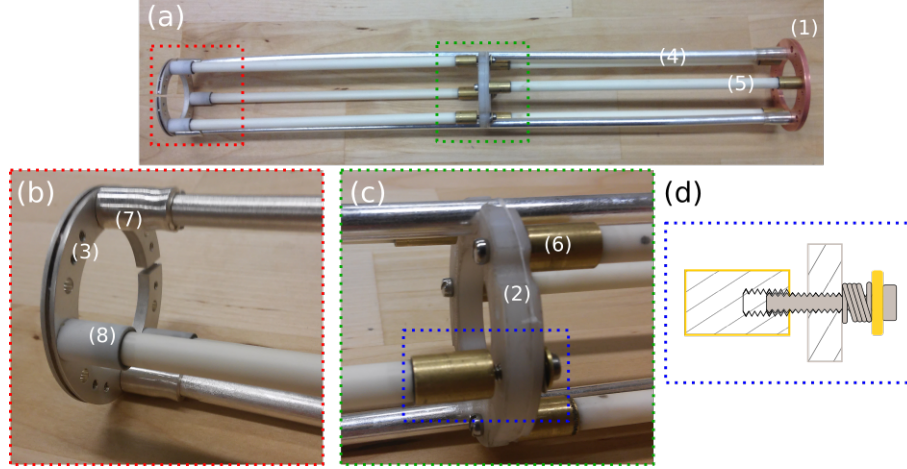


Figure 7.3: **Cold-finger design:** (a) Photo of the cold-finger. A copper disk (1) is attached to the MXC plate, and annealed silver rods (4) provide thermalization to the silver sample mounting disk (3). Alumina rods (5) provide structural support. Close up views of (b) the sample mounting disk and (c) the teflon intermediate disk. (d) Mechanism at the intermediate disk for compensating for the discrepancy between the thermal expansion coefficients of alumina and silver upon cooling to cryogenic temperatures.

and lattice defects in the metal and decrease the low-temperature resistivity of the silver.⁶ Silver is a convenient choice of metal, since it doesn't superconduct and has a small nuclear heat capacity at dilution refrigerator temperatures [89]. Since the annealing process removes grain boundaries, the metal becomes very soft and machining becomes difficult. We therefore use un-annealed silver (Fig 7.3, (7)) crimped onto the ends of the rods, which are tapped and screwed onto the MXC mounting disk and sample mounting disks ((Fig 7.3, (1) and (3) respectively.)

To prevent eddy-current heating when we sweep the magnetic field, we would like to minimize the amount of metal in thermal contact to the sample and, when metal is unavoidable, prevent large conducting loops that pick up large EMF's. The sample mounting disk, which

⁶The quality of a metal as a low-temperature conductor is generally quantified by the residual resistivity ratio (RRR) defined as $\rho[300\text{ K}]/\rho[4\text{ K}]$. This is generally also taken as a quantifier of the purity of the metal, since at room temperature resistivity is dominated by electron-phonon scattering, while at liquid helium temperatures resistivity is generally dominated by impurity/grain boundary scattering.

must be conducting to thermally anchor the sample, is machined with a notch to break such a loop. For structural support, we use alumina ceramic rods (Fig 7.3, (5)), capped off with tapped brass connectors (Fig 7.3, (6)) since alumina is prohibitive to machine. On the sample mounting disk, near the magnetic field maximum, we instead use caps made of Macor, a machinable ceramic (Fig 7.3, (8)) to connect the sample mounting silver flange.

When designing structures at cryogenic temperatures, it is critical to take thermal contractions into account. A major problem with the cold-finger scheme presented here is that the thermalizing silver has a much higher coefficient of thermal expansion than the structural alumina. If uncompensated, the relative contraction of the silver rods compared to the alumina rods would likely deform or destroy the structure. To circumvent this, we implement a spring-loaded stress relief mechanism⁷ (Fig 7.3(d)), where we intentionally design a gap between one set of alumina rods and the mid-finger teflon disk (Fig 7.3, (2)). At room temperature, a compressed spring on the other other side of the disk provides structural stability to the rods; as we cool and the silver contracts relative to the alumina the spring decompresses, decreasing the overall length of the structural component. This mechanism has the added benefit that, for any reasonable experimental mass, the resonant oscillation frequency of such a spring/mass system is far detuned from the pulse tube frequency (~ 1 Hz) giving the cold finger intrinsic vibration isolation.

The sample stage, shown in Fig. 7.4 is also machined out of silver, with slots cut into the large-radius areas to prevent eddy current heating. The sample stage stage is designed such that 16 DC measurement lines and 4 RF measurement ports may be used in an experiment. Each DC measurement line is connected to cryostat ground via a 470 pF mylar

⁷We actually realized this problem after designing the cold finger, and this mechanism was designed on the spot by Bill Pratt when we asked him for advice. I think he may be a wizard.

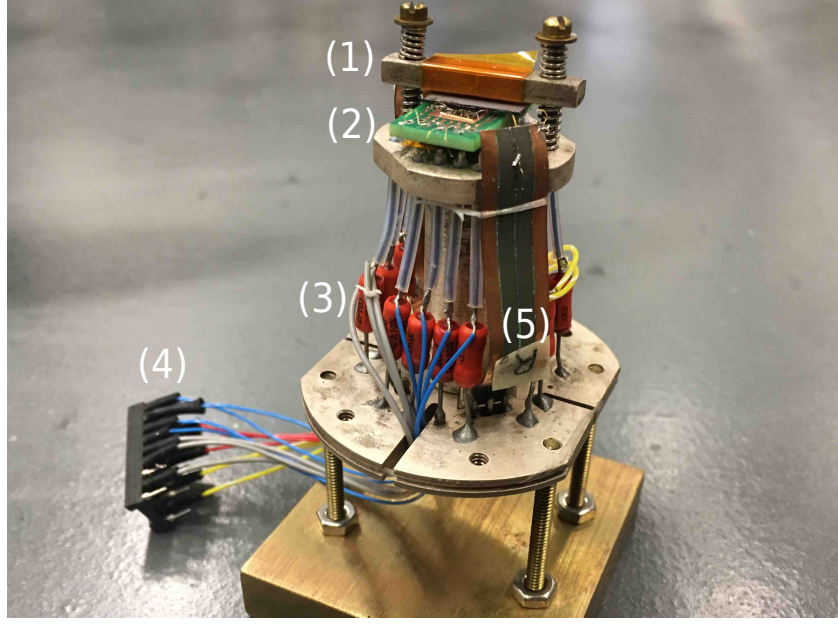


Figure 7.4: **Acoustoelectrics sample stage**: Picture of the sample stage. (1) We use a spring loaded mechanism to clamp the SAW device down to the graphene device, which is mounted on a custom printed circuit board (2) for wire-up. (3) Thin film mylar capacitors (red components) provide a high frequency short to ground as well as thermalization for the electrons. (4) The leads are wired up to an 18-pin connector, which connects to the standard wiring in the fridge. (5) SMA panel-mount connectors at the top of the sample stage (bottom in this picture) are connected to a strip-line resonator on flexible kapton substrate, which provides the RF signal to the SAW device.

thin-film capacitor (Quest Components, part No. 192P471X9200.) These capacitors serve two roles: (1) the wrapped up mylar thin-film provides a large surface area contact between the electrical leads (and thus the 2DES) and the cold cryostat. This reduced the thermal resistance between the 2DES and the cryostat while maintaining an infinite DC electrical resistance, alleviating the need for more extreme measures to cool the 2DES such as a helium-3 immersion cell [108]. (2) Along with an in-line 10 k Ω metal-film resistor on each measurement line, the capacitors form a lumped-element low pass filter that attenuates noise above $f_{LP} \sim (RC)^{-1} \approx 200$ kHz.

The RF ports are fed in via SMA surface mount connectors press fitted and soldered into

the top of the stage. To run the RF signal from the top of the stage to the device mounted on the bottom, we used a strip of double-sided copper-clad Kapton, with a micro-stip line etched on one side such that $Z_{strip} \sim 50 \Omega$. From there, we used thin gold wires to solder from the SAW transducers to the RF strip-lines. This was sufficient for the ~ 300 MHz SAW excitation signals used in this experiment, however we never tried to quantify losses or parasitic resonances. At higher frequencies, where both of these problems are exacerbated, it would likely be simpler and more efficient to run a coaxial cable directly to the level of the sample head and connect to a surface mount connector on a printed circuit board with prefabricated strip-lines.

7.3.2 Experimental setup

The graphene device used in this experiment was fabricated by Erik Henriksen's group at Washington University at St. Louis. The monolayer graphene sample was mechanically exfoliated onto a degenerately doped silicon wafer with a 300 nm SiO_2 surface layer. Electrical leads for conventional and acoustoelectric transport were fabricated such that they crossed the entire width of the graphene sample perpendicular to the SAW propagation direction as shown in Fig. 7.5(a). These leads were defined using standard lithographic techniques, and had a width of $3 \mu\text{m}$. Metallization of these electrical contacts was achieved by thermal evaporation of 4 nm of chromium followed by 70 nm of gold. A voltage applied to the entire silicon substrate was used to gate-control the graphene charge carrier density *in situ* as shown in Fig. 7.5(b).

Silicon is not piezoelectric and hence SAW measurements cannot be done using a silicon substrate alone. Therefore, we employed a flip-chip configuration, where a pair of interdigitated transducers (IDTs) for launching SAWs were fabricated onto a separate piezoelectric

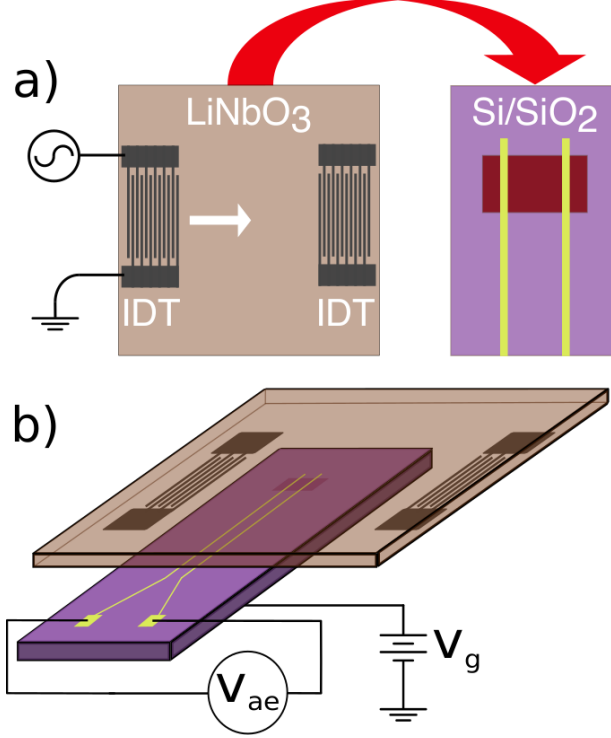


Figure 7.5: **Graphene acoustoelectrics device schematic:** Schematic of the experimental setup. (a) The flip-chip configuration for the acoustoelectric measurements consists of two devices: a SAW delay line fabricated onto a lithium niobate “black” substrate, and an exfoliated graphene device on an oxidized silicon chip. A high-frequency signal applied to one of the interdigitated transducers forming the delay line launches a SAW in the direction indicated by the white arrow. (b) The SAW delay line is flipped onto and mechanically clamped to the graphene device, as described in the text. In this configuration the electric field propagating in tandem with the SAW evanescently couples to the graphene device, while the silicon back gate may be used to control the graphene charge carrier density *in situ*. The resulting acoustoelectric voltage is measured between two electrical contacts that extended across the entire width of the graphene sample.

substrate that was then mechanically pressed onto the silicon substrate hosting the graphene device, as shown in Fig. 7.5(b). In our device, aluminum IDTs were fabricated in the form of a SAW delay line onto a single-crystalline chip of 128° Y-X lithium niobate “black” via conventional photolithography. The IDT pair had a center-to-center spacing of 7 mm with each transducer composed of 40 pairs of $3\ \mu\text{m}$ wide interdigitated electrode fingers.

A SAW is launched along the surface of the lithium niobate by applying an alternating

voltage between the two electrodes of an IDT at the fundamental frequency of the transducers, $f = v/\lambda \simeq 332$ MHz. This frequency corresponds to a SAW wavelength $\lambda \simeq 12 \mu\text{m}$. The electric field accompanying the SAW evanescently extends above the surface of the lithium niobate, meaning the SAW may couple to a 2DES at a height $< \lambda$ above the lithium niobate [233, 234]. To effectively couple the SAW to the charge carriers in the graphene, care was taken to ensure the absence of large debris on either substrate that could become lodged between them after being flipped into contact. The two chips were mounted in a custom spring-loaded sample holder to maintain their contact at low temperature and to immobilize their relative lateral motion. Previous studies [233, 202, 235, 236] of flip-chip assemblies similar to ours have estimated air gaps substantially smaller than the SAW wavelength in our experiment. Therefore, while we are not able to directly measure the air gap in our device, we expect it to be sufficiently small to enable acoustoelectric coupling.

With this experimental setup, low-temperature acoustoelectric measurements were made between pairs of leads on the graphene by applying an amplitude modulated signal to one SAW transducer and detecting the correspondingly generated acoustoelectric voltage V_{ae} using standard ac lock-in techniques. These measurements correspond to the open configuration in which the total acoustoelectric current density $j = 0$ (see Eqn. 7.12) and V_{ae} arises from the charge displacement produced by the SAW in the direction of its propagation. We note that the electrical leads used for detected V_{ae} also allow us to perform conventional lock-in based low-frequency (13 Hz) transport measurements to characterize the graphene sample and provide a point of reference when interpreting our acoustoelectric data. Finally, all data were taken in zero magnetic field and $T \simeq 3.2$ K.

7.3.3 SAW delay line characterization

The frequency response of the SAW transducers in our flip-chip device was characterized using an Agilent N5230A vector network analyzer. Fig. 7.6(a) shows the measured reflection coefficient S_{11} of the SAW device as a function of frequency. The resonance in the reflected power observed at $\simeq 337$ MHz is associated with the generation of SAWs in the LiNbO₃. The slight shift in frequency relative to the expected SAW resonance at 332 MHz is attributable to an increase in the elastic moduli of LiNbO₃ and mechanical strain in the flip-chip device upon cooling to cryogenic temperatures⁸. The primary SAW resonance shown in Fig. 7.6(a) also exhibits superimposed oscillations as a function of frequency. Examining the inverse Fourier transform of S_{11} , shown in the inset of Fig. 7.6(a), reveals a series of decaying peaks spaced by $3.4 \pm 0.1 \mu\text{s}$. This time scale corresponds to a SAW propagation distance of 13.3 ± 0.2 mm, which is approximately twice the spacing of the two transducers in our device. Because the transducers also act to reflect surface acoustic waves, a launched SAW will propagate across the chip, be reflected by the opposing transducer, and propagate back to the launching transducer, where its co-propagating electric field will interfere with the reflected electric signal from the launching transducer as measured by the network analyzer. We therefore attribute these superimposed oscillations in S_{11} to the interference between the signal reflected from the launching transducer and the electric field associated with SAWs that have propagated to and from the opposing transducer.

⁸We note that mass loading alone from the SiO₂/Si substrate would tend to reduce the resonant response of the SAW delay line. However, since the resonance frequency is shifted slightly higher by ~ 5 MHz, we conclude that stiffening of the lithium niobate crystal due to strain and cooling to low temperatures dominate mass loading by the SiO₂/Si substrate.

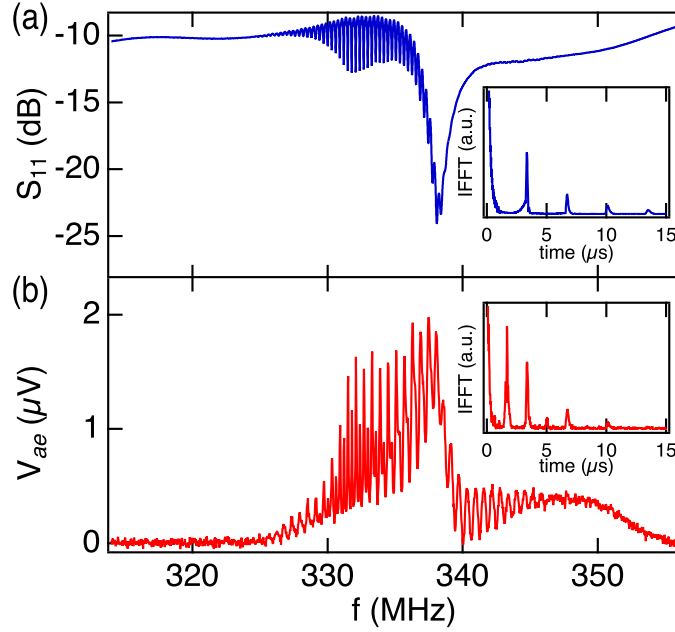


Figure 7.6: **SAW reflection and acoustoelectric voltage vs. drive frequency:** (a) Frequency dependence of the reflection coefficient (S_{11}) of the SAW delay line. Inset: inverse fast Fourier transform (IFFT) of S_{11} data. (b) Measured acoustoelectric voltage in the graphene sample as a function of frequency. This measurement was conducted at 3.2 K with the sample back gate grounded. To excite the SAW, 0 dBm of microwave power was applied at the top of the cryostat. Inset: inverse Fourier transform of the acoustoelectric signal.

7.3.4 Graphene acoustoelectrics

Fig. 7.6(b) shows the acoustoelectric voltage measured between two leads $29\ \mu\text{m}$ apart⁹ on the graphene device as a function of frequency. The resonance in V_{ae} coincident with the generation of SAWs centered at $\simeq 337\ \text{MHz}$ is associated with the acoustoelectric transport of charge in the graphene sample. For this data the silicon back gate was held at ground potential and, as we will describe below, the sample displays p-type conduction typical of graphene devices on SiO_2 exposed to air and polymer resist. The sign of the induced voltage is consistent with conduction by holes.

The frequency response of V_{ae} shows oscillations superimposed onto the main peak similar

⁹Consistent behavior was observed using the other leads on the device.

to those observed in S_{11} . These oscillations are attributable to the modulation in the SAW amplitude caused by interference between the primary SAW and reflected waves.[221] To quantitatively understand the nature of these oscillations, we examine the inverse Fourier transform of V_{ae} , which is shown in the inset of Fig. 7.6(b). Similar to the behavior of S_{11} we observe a series of decaying peaks spaced by $\simeq 3.4 \mu\text{s}$. We attribute these peaks to the interference of the primary SAW with forward propagating SAWs that have been reflected by the far transducer and then again by the original transducer. Additionally, we observe a number of peaks at odd multiples of $1.7 \pm 0.1 \mu\text{s}$, which correspond to the propagation distance between the two transducers. We attribute these peaks to modulation in the SAW amplitude caused by counter propagating waves reflected by the opposing transducer.

To verify that the measured voltage is associated with the charge carriers in the graphene, we applied a voltage V_g to the silicon back gate to tune the graphene charge carrier type from holes to electrons. Fig. 7.7(a) shows the acoustoelectric voltage as a function of frequency for two back gate voltages on either side of the Dirac peak (see Fig. 7.8(a)). With the back gate grounded ($V_g = 0 \text{ V}$, red trace Fig. 7.7(a)), the graphene exhibits conduction due to holes and the corresponding acoustoelectric signal is positive. By increasing the V_g to $+60 \text{ V}$, electrons become the predominant charge carrier in the graphene and, as expected, the acoustoelectric voltage correspondingly reverses sign (blue trace in Fig. 7.7(a)). In Fig. 7.8 we show the full gate- and frequency-dependent map of V_{ae} for the device along with the corresponding gate-dependent two-point resistance of the graphene. In the bottom panel of Fig. 7.8 we show a constant frequency linecut of V_{ae} at $f = 337.45 \text{ MHz}$. As expected, the sign of acoustoelectric signal changes upon tuning through the Dirac peak at $V_g \simeq 28 \text{ V}$, validating that V_{ae} is a result of acoustoelectrically induced charge transport in the graphene.

Moreover, in the open configuration of our measurement (where $j = 0$) Eqn. 7.12 predicts

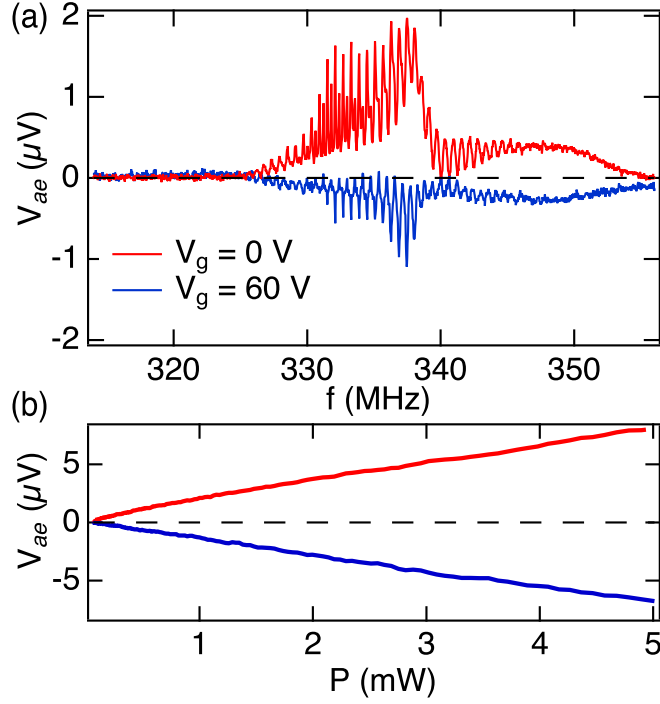


Figure 7.7: **Acoustoelectric effect: charge carrier polarity dependence:** (a) Acoustoelectric voltage versus frequency for two gate voltages: $V_g = 0$ V (hole-doped graphene) and $V_g = 60$ V (electron doped graphene) at 3.2 K. For these measurements 0 dBm of microwave power was used to excite SAWs. (b) Acoustoelectric voltage as a function of power applied to the SAW circuit, for the same back gate voltages as in panel (a). This data was taken at a fixed frequency $f = 337.45$ MHz corresponding to the maximum in the response of the SAW transducer.

that V_{ae} should be proportional to the SAW intensity I . Therefore, to further verify the acoustoelectric origin of the signal, we measured V_{ae} as a function of SAW power at a fixed frequency of $f = 337.45$ MHz, which corresponds to the peak in SAW resonant response. The results are plotted in Fig. 3(b). For both hole and electron doping we find that V_{ae} is linear in the applied SAW power, consistent with Eqn. 7.12 and with previous acoustoelectric measurements in graphene [227, 229, 230, 231].

In the vicinity of the Dirac peak we observe a marked reduction in the magnitude of

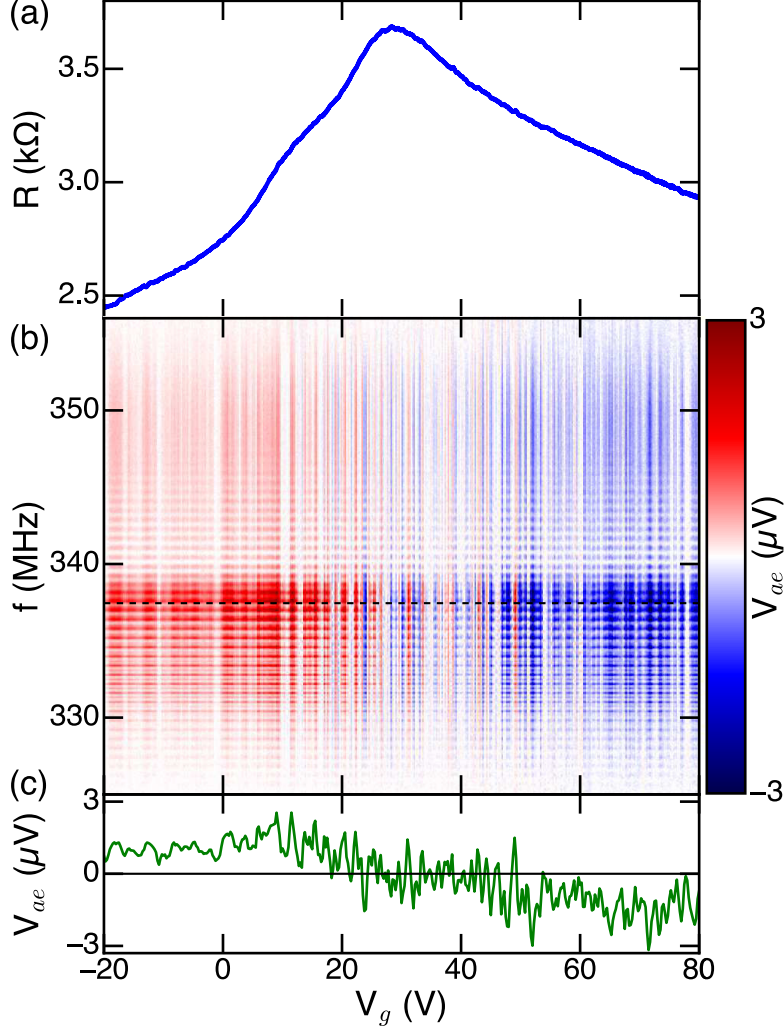


Figure 7.8: **Gate-tunable acoustoelectric effect:** (a) Low-frequency two-point resistance of the graphene sample as a function of back gate voltage. (b) Acoustoelectric voltage signal (taken with 0 dBm of SAW power), plotted as a function of both back gate voltage and SAW frequency. At gate voltages where the graphene is heavily electron (hole) doped, V_{ae} is consistently negative (positive). (c) Constant frequency linecut of V_{ae} at $f = 337.45$ MHz, indicated by the dashed line in panel (b).

$V_{ae} \rightarrow 0$ over a well-defined region in gate voltage. This region is emphasized in Fig. 7.9(a), where we performed measurements with increasing levels of SAW excitation power to enhance the magnitude of the acoustoelectric signal. For comparison in Fig. 7.9(b) we show low frequency transport over the same range of gate voltage while applying a SAW driving signal of the same power. Over a range of ~ 20 V, as indicated by the dashed vertical lines, we find

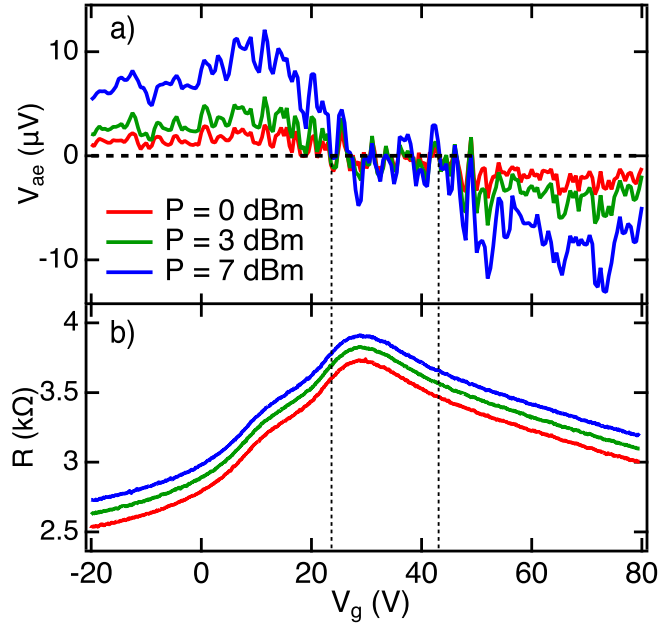


Figure 7.9: **Power dependence of the Gate-tunable acoustoelectric effect:** (a) Acoustoelectric signal versus back gate voltage at a SAW frequency $f = 337.45$ MHz for increasing levels of SAW power. A well-defined region where $V_{ae} = 0$ is observed around charge neutrality and is associated with heterogenous charge disorder as discussed in the text. (b) Corresponding low-frequency two-point resistance of the graphene sample as a function of back gate voltage for the same values of SAW power as shown in (a). The two-point resistance traces in have been offset by $100 \text{ } \Omega$ relatively to each other for clarity.

that the average value of $V_{ae} \simeq 0$. This range of gate voltage corresponds to a charge carrier density range of approximately $1.4 \times 10^{12} \text{ cm}^{-2}$. This region of suppressed V_{ae} is likely associated with the formation of heterogeneously doped electron-hole “puddles”, which are known to exist close to charge neutrality due to substrate contamination and impurities [237, 238] but only manifest indirectly in transport measurements as a broadening of the Dirac peak. In this charge disordered regime one would expect that competing acoustoelectric signals from roughly equal distributions of electrons and holes to cancel each other and result in an approximately net zero value of V_{ae} as observed in our measurements. From the low frequency transport we determine that the field effect mobility of the graphene

is $\mu \simeq 320 \text{ cm}^2/\text{Vs}$ on the hole-doped side of the Dirac peak and $\simeq 110 \text{ cm}^2/\text{Vs}$ on the electron-doped side. While these values are relatively low for an exfoliated device they are, however, two orders of magnitude higher than those reported in ion-gated devices used in recent acoustoelectric measurements [231, 239]. Moreover, the fact that the mobility of holes in our sample is roughly three times larger than that of electrons naturally explains the asymmetric position of the region over which $V_{ae} \simeq 0$ relative to the Dirac peak. Finally we note that a reduction in SAW driven *current* was observed near charge neutrality in recent measurements of CVD graphene at room temperature [231, 239, 240] using ion gel gating. In these measurements the authors suggest similarly that electron-hole puddling is responsible. Our measurements using a different experimental setup and an exfoliated device at low temperatures are consistent with this interpretation and further highlight that acoustoelectric methods provide complementary methods for investigating 2DES and reveal features of the underlying phenomena not directly manifest in conventional transport.

7.4 Future work

There are several directions one could imagine taking the concepts developed in this experiment. We present several of them here:

7.4.1 Graphene acoustoelectrics in a magnetic field

By depositing leads across the graphene and perpendicular to the direction of SAW propagation, this experiment was designed to measure acoustoelectric current flowing in the direction of propagating SAW waves, i.e. when the conductivity can be written as a *scalar*. As is evident from Eqn. 7.10, when a magnetic field is introduced and the conductivity becomes a

tensor, the acoustoelectric current need not flow in the direction of SAW propagation [213]. In-field acoustoelectric measurements have been done in GaAs [212, 204, 213] in a more traditional Hall bar geometry. It would be fairly straightforward to fabricate a graphene device with Hall bar geometry, which could potentially allow for these measurements to be extended into the quantum Hall regime.

7.4.2 $\Delta v/v$ measurements

A modern vector network analyzer (VNA) is likely capable of measuring the phase shift of a microwave signal with a higher precision than the phase-locked loop setup typically employed in $\Delta v/v$ measurements [168, 206]. Using a VNA, it might be possible to do $\Delta v/v$ measurements on a piece of exfoliated graphene. The transducers would need to be made small, such that the width of the SAW propagation path was comparable to the size of the exfoliated device (~ 10 s of μm), and the transducers would have to be put relatively close to one another such that the exfoliated device takes up a significant portion of the SAW propagation path. For example, for a SAW delay line with two $20\ \mu\text{m}$ wide transducers spaced by $100\ \mu\text{m}$, a $20 \times 10\ \mu\text{m}$ piece of exfoliated graphene, an eminently reasonable device size, would cover $\approx 1/10$ of the SAW propagation path.

In this proposed experiment, if the microwave environment isn't carefully engineered to avoid spurious modes, the close SAW transducers would likely have a large free-space coupling that completely washes out any SAW coupling. This is, however, a manageable problem, as displayed by several quantum acoustics experiments with similar SAW propagation path sizes with a clear SAW propagation signal [7, 17, 174]. We refer the reader to Ref. [241] for an introduction on engineering experiments to avoid spurious microwave modes.

7.4.3 Measurements of 2DESs using SAW resonators

An interesting possibility¹⁰ for using SAWs in the quantum Hall regime was inspired by Ref. [242] where composite fermion cyclotron resonances were resistively detected when they became commensurate with a piezoelectrically generated periodic electric field in GaAs. In this paper, the periodic strain field that generates the electric field was created by stripes of photoresist on the surface of the substrate, which upon cooling to the base temperature of the cryostat contracted relative to the GaAs substrate. It may be possible to do a similar resistively detected cyclotron resonance experiment by placing a graphene device in a SAW resonator. Exciting the SAW resonator at its resonant frequency would set up a standing wave, which introduces a length scale (the SAW wavelength λ) which may modulate the transport properties when the radius of the cyclotron orbit of a composite fermion is commensurate with λ [185, 213].

A particularly tantalizing idea is that if the Bragg mirrors are far enough apart, the free spectral range of the SAW cavity may easily be engineered to be small enough such that the cavity hosts several resonant modes with slightly different frequencies/wave vectors [169]. Each of these resonant modes would subsequently be commensurate with slightly different composite fermion cyclotron resonances. This would, in principle, allow for one to investigate cyclotron resonances as a function of frequency on a single sample (albeit over a small frequency range set by the stop-band of the mirrors, see § 5.3.5) in a resistively detected manner similar to Ref. [242].

Alternatively, one may be able to infer $\Delta v/v$ and Γ from shifts in the real and imaginary ($\propto$ the linewidth of the resonance) parts of the SAW resonant frequency respectively. A similar technique was employed in Ref. [243] to study the bath of two-level systems near

¹⁰Thank's to Leo Li for this idea.

the surface of a piezoelectric substrate. Given other topics in this thesis, one could imagine a situation where these measurements could be enhanced by the incorporation of coherent quantum circuits coupled such a SAW resonator. For example, fine shifts in the resonant frequency of the cavity could be detected with high precision by measuring the phase accumulated by a single phonon in the resonator relative to that of a superconducting qubit coupled to the resonator.

APPENDICES

APPENDIX A

A note on the rotating wave approximation

At several points in this thesis, we have invoked something called the “rotating wave approximation” or RWA. When considering Hamiltonians such as Eqn. 3.13 that contain a coherent drive term, when we make the transformation to the frame rotating with the drive we will in general have terms in the Hamiltonian that are a *sum* of two frequencies and terms that are a *difference* of two frequencies. The RWA may be invoked when there is a separation of scale between the oscillation frequency of the sum (“counter-rotating”) terms and characteristic rates in the rest of the Hamiltonian: if the counter rotating terms oscillate quickly over the timescale of the rest of the dynamics, the net effect of these term will average out over many periods and we can drop them.

For example, consider the derivation of Eqn. 3.17: before invoking the RWA, the Rabi Hamiltonian in the frame rotating is

$$\tilde{\mathbf{H}}_{\text{Rabi}} = -\hbar \frac{\omega_{01} - \omega_d}{2} \boldsymbol{\sigma}_z - \frac{A}{2} \left(\cos(\phi) \boldsymbol{\sigma}_x + \sin(\phi) \boldsymbol{\sigma}_y + e^{-2i\omega_d t} \boldsymbol{\sigma}_+ + e^{2i\omega_d t} \boldsymbol{\sigma}_- \right) \quad (\text{A.1})$$

The time dependence of the $\boldsymbol{\sigma}_x$ and $\boldsymbol{\sigma}_y$ terms, the “co-rotating” terms, has been can-

celed out, while the counter-rotating terms (σ_+ and σ_-) have picked up time dependencies $\propto e^{\pm 2i\omega_d t}$. The main argument of the RWA is that, from the Schrödinger equation $i\hbar\partial_t\psi = \mathbf{H}\psi$ the characteristic rates at which the state vector will change are $|\omega_{01} - \omega_d|$ and $A/\hbar$. If $\omega_d \gg \max[|\omega_{01} - \omega_d|, A/\hbar]$, the counter-rotating terms will oscillate quickly compared to the evolution of the state vector, and the effect of these terms will average out to zero.

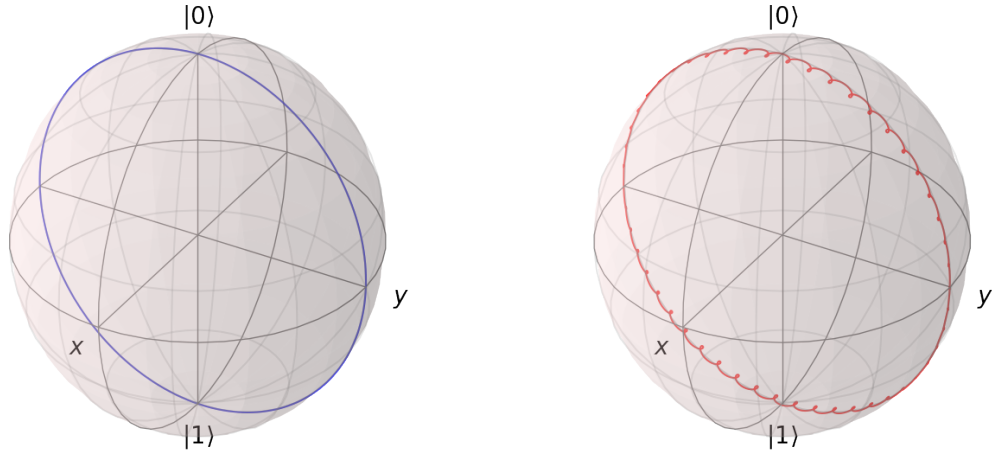


Figure A.1: **The rotating wave approximation:** (a) State evolution after invoking the rotating wave approximation (RWA). Here, we simulate a qubit subject to a resonant Rabi drive ($\omega_q = \omega_d$), with $\Delta = 0$ such that the state vector rotates around the x -axis of the Bloch sphere. We have chosen $\omega_q/(2\pi) = 5$ GHz and $A/(2\pi) = 0.1$ GHz. (b) The same state evolution without invoking RWA. Small oscillations in the path are averaged out over time-scales $\approx 1/A$, and the state-vector largely takes the same path as if we had invoked the RWA. Figures made using the QuTiP package [102].

A visualization of this is given in Fig. A.1, where we have numerically solved the Schrödinger equation for the Hamiltonian Eqn. A.1 with (b) and without (a) the counter-rotating terms and plotted the corresponding trajectories on the Bloch sphere. For this simulation, we've taken resonant qubit and drive frequencies $\omega_{01}/(2\pi) = \omega_d/(2\pi) = 5$ GHz and drive amplitude $A/(2\pi) = 100$ MHz. Even for this relatively strong Rabi drive (corre-

sponding to a π -pulse duration of ~ 5 ns), including the counter-rotating terms does little to effect the overall path of the Bloch-vector.

The RWA is also employed several times in Chapter 2: in those cases, the best way to think about the RWA is in the Heisenberg picture, where the operators are time dependent. The raising and lowering operators of the unperturbed harmonic oscillator have time dependencies $\mathbf{a}^\dagger(t) = \mathbf{a}^\dagger(0)e^{-i\omega t}$ and $\mathbf{a}(t) = \mathbf{a}(0)e^{i\omega t}$ (exercise: prove this.) Therefore, when deriving Eqn. 2.27, terms such as $\mathbf{a}^\dagger\mathbf{a}^\dagger\mathbf{a}^\dagger\mathbf{a}^\dagger$ will oscillate quickly compared to the dynamics of the rest of the Hamiltonian, and may be ignored. When deriving the JCH (Eqn. 2.41), dropping terms such as $\mathbf{a}_q^\dagger\mathbf{a}_c^\dagger$ is only strictly valid when $|\omega_{01} - \omega_c| \ll |\omega_{01} + \omega_c|$, i.e. when the qubit and the cavity aren't too far detuned from each other. However, the JCH and the dispersive regime approximation tend to be valid for detunings that are an appreciable fraction of the qubit frequency ($\Delta/(2\pi) \approx 2 - 3$ GHz is very common.)

APPENDIX B

Other cQED experimental techniques

In this appendix, we summarize a couple more advanced experimental protocols we have employed while experimenting on superconducting circuits.

Excited state population

In an ideal scenario, the baseline excited state P_1 population of a qubit would follow a Boltzmann distribution. If this were the case $\omega_{01} \sim 5 \text{ GHz} \sim 250 \text{ mK}$, so at typical dilution refrigerator operating temperatures ($\sim 10 - 20 \text{ mK}$) P_1 would be exponentially suppressed to a negligible level¹. However, spurious excitation modes (attributed mainly to non-equilibrium quasiparticles [141, 154, 142, 42]) may populate the excited state to well above the level expected by Maxwell-Boltzmann statistics. As such, it is useful to measure the residual excited state population.

We accomplish this using a method developed in Ref. [158] and further optimized in Ref. [154], where we measure the ratio of the amplitude of two Rabi drives: one where we Rabi flop the residual excited state population $P_{1,R}$, and one where we Rabi flop the residual ground state population $P_{0,R}$. To do these two measurements independently, we need a third energy level that starts off with negligible population to act as the target of the Rabi drive. Transmon qubits have small anharmonicity, meaning that the $|2\rangle$ state can fulfill this role

¹This is equivalent to saying that $\Gamma_{\uparrow}$, the transition rate from $|0\rangle$ to $|1\rangle$, is exponentially suppressed.

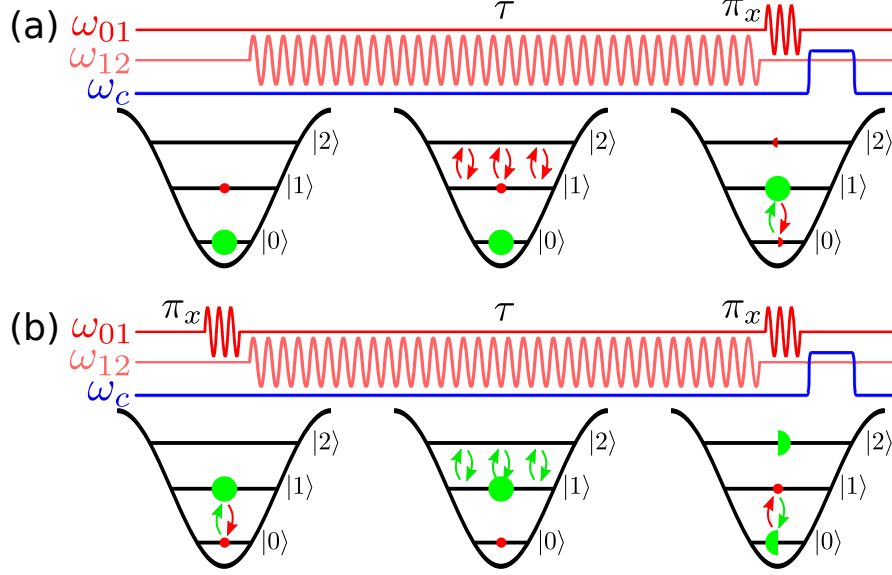


Figure B.1: **Excited state population measurement:** To measure the residual excited state population, we measure the relative amplitude of two Rabi drives, one (a) where we Rabi drive the residual excited state population ($P_{1,R}$, red circles) between the $|1\rangle$ and $|2\rangle$ states, and one (b) where we first invert the population of $|0\rangle$ and $|1\rangle$ such that we Rabi drive the residual *ground* state population ($P_{0,R}$, green circles) between the $|1\rangle$ and $|2\rangle$ states. At the end of each measurement, we invert the population of $|0\rangle$ and $|1\rangle$ and measure.

without too much experimental effort.

The first measurement simply consists of driving the qubit at ω_{12} for a variable time τ , which Rabi flops the population $P_{1,R}$ between $|1\rangle$ and $|2\rangle$, see Fig. B.1(a). To measure, we apply a π pulse at ω_{01} , and measure, repeating the measurement while varying τ to build up $P_0(\tau)$, which should be sinusoidal as predicted by Eqn. 3.14. We fit $P_0(\tau)$ to a sinusoid, and extract the amplitude A_0 .

We then repeat the measurement, except this time before the ω_{12} Rabi drive we apply a π pulse at ω_{01} , inverting $P_{0,R}$ and $P_{1,R}$. Now, the drive Rabi flops $P_{0,R}$, the residual ground state population, between $|1\rangle$ and $|2\rangle$, see Fig. B.1(b). We repeat the rest of the process, building up $P_0(\tau)$ and fitting it to a sinusoid, extracting the amplitude A_1 . Then, if the

qubit $|0\rangle$ and $|1\rangle$ states are the only residually populated states², $P_{1,R}$ is simply given by

$$P_{1,R} = \frac{A_0}{A_1 + A_0} \quad (\text{B.1})$$

It was noted in Ref. [154] that since the amplitudes of the sinusoids are really the only fit parameters we care about in this measurement. every point of the Rabi drive we measure that isn't at a maximum or minimum of the sine ($\tau = n\pi/\Omega_R$, where n is an integer) is functionally dead weight. Thus, a more efficient measurement is to replace the variable time τ with a variable number of π pulses between $|1\rangle$ and $|2\rangle$ (i.e., a Rabi drive where we only sample at $\tau = n\pi/\Omega_R$.) With this modified measurement scheme, we can significantly expedite the data taking process, and average over a far higher number of measurements to record $P_{1,R}$ at much higher precision (we have measured $P_{1,R}$ down to $\sim 0.5\%$.)

Spin-locking noise spectroscopy

As outlined in § 4.2 decay (T_1) measurements are sensitive to the noise spectral density at $\pm\omega_{01}$, while dephasing (T_2) measurements sample a wide range of the noise spectral density. Often times, we want to reconstruct the noise spectral density at frequencies much lower than ω_{01} , from mHz to MHz, since noise at these frequencies contribute to qubit dephasing [105, 106, 62], and since measuring the characteristics of the noise PSD is critical for feeding into theoretical noise models and engineering experimental setups that minimize noise [134, 104, 128, 88, 244]. However, these frequencies are not accessible with simple T_1 measurements, the wide range of noise spectrum sampled in a dephasing mechanism makes it difficult to deconvolute the spectrum from measurements. What we would like is

²This is a good approximation, given that we typically measure $P_{1,R} \approx 0.5 - 3\%$.

a measurement that samples a sharp frequency range (like a decay measurement) but can access these frequencies. Spin-locking measurements [104, 245, 62, 88] are one method of doing this.

Experimentally, we accomplish spin-locking measurements with something similar to a T_2 measurement: we start off with the qubit in the ground state, and apply a $-\pi_y/2$ pulse to rotate the state vector on the onto the x axis of the Bloch sphere (see Fig. B.2(b).) Then, instead of simply waiting time τ before the next $\pi/2$ pulse, during this time we apply a on-resonant ($\omega_{01} - \omega_d = 0$) microwave drive along the x axis ($\phi = 0$ in Eqn. (3.17).) These conditions render the Hamiltonian in the rotating frame to be

$$\tilde{H} = -\frac{\hbar\Omega_R}{2}\sigma_x \quad (\text{B.2})$$

The state vector we have prepared is in an eignestate of this rotating frame Hamiltonian, and will thus only undergo non-deterministic evolution (similar to a T_1 measurement in the lab frame.), decaying exponentially at rate $\Gamma_{1\rho}$ as is depicted in Fig. B.2(b). After time τ , we apply another $-\pi_y/2$ pulse to align the state vector with the z axis of the Bloch sphere and perform a projective measurement. As above, we repeat this process varying τ , and repeat this series a sufficient number of times to accurately measure $P_1(\tau)$. Since during evolution our state vector is in an eigenstate of the rotating frame Hamiltonian, we expect to only measure non-deterministic exponential state decay. We thus fit $P_1(\tau)$ to a decaying exponential to extract $\Gamma_{1\rho}$. This decay rate is given by [245]

$$\Gamma_{1,\rho} = \frac{1}{2}\Gamma_1 + \Gamma_\nu \quad (\text{B.3})$$

where Γ_ν is proportional to the noise spectral density at Ω_R . To extract Γ_ν , we need to

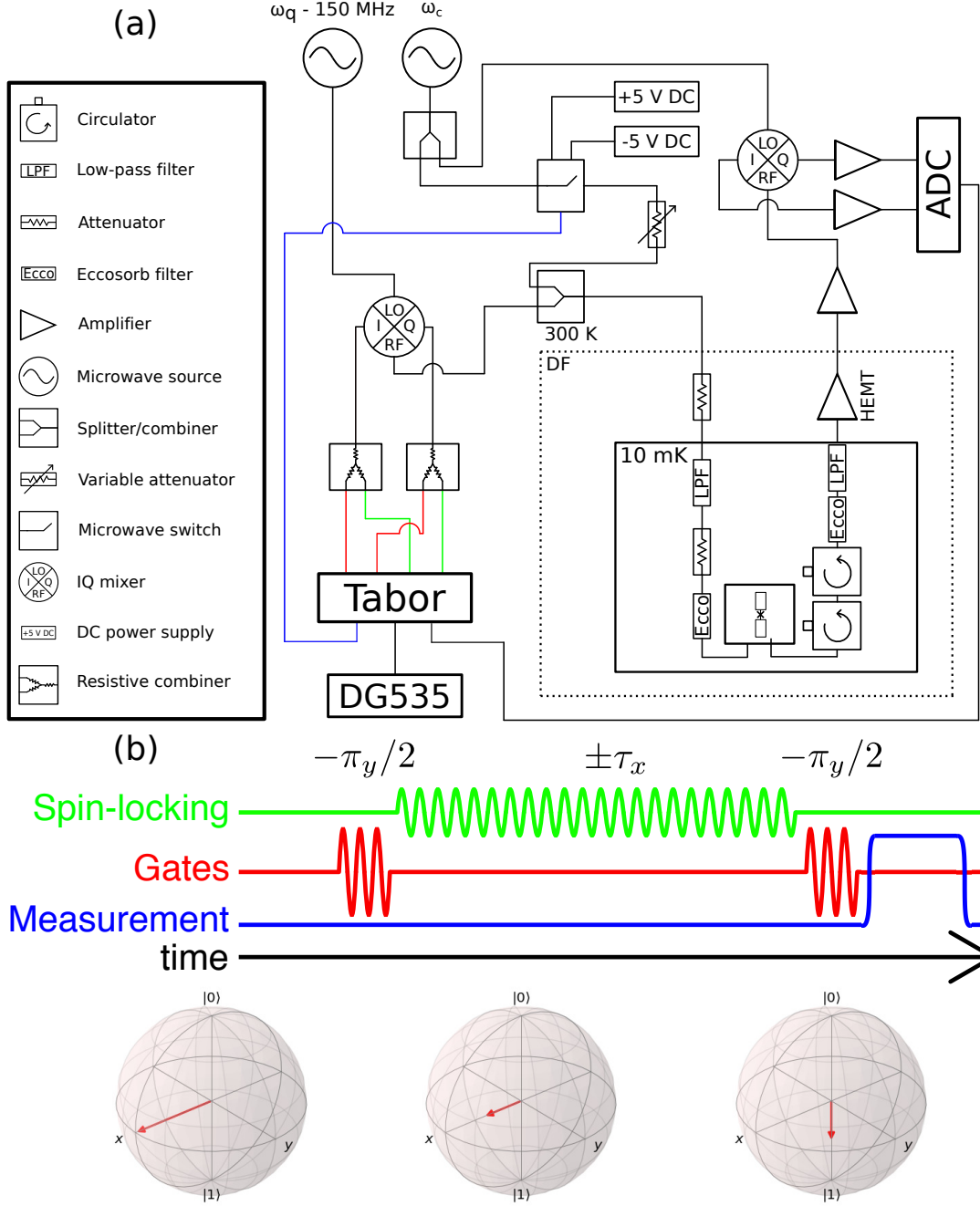


Figure B.2: **Spin-locking setup and measurement:** (a) Modified measurement setup for spin-locking measurements. To vary the amplitude of the locking Rabi drive (green) while keeping the control pulse amplitudes (red) constant, we control them using different pairs of channels on the Tabor. The cavity pulse is then generated using a microwave switch actuated by a voltage pulse from a Tabor marker channel (blue.) (b) Pulse sequence and state evolution. We start with the qubit in the ground state, and then apply a $-\pi_y/2$ pulse to prepare the state vector along the x axis. We then “lock” the qubit along the x axis by turning on a Rabi drive, which puts the state in an eigenvector of the rotating frame Hamiltonian. After time τ , we perform another $-\pi_y/2$ pulse to put our Bloch vector along the measurement axis and measure.

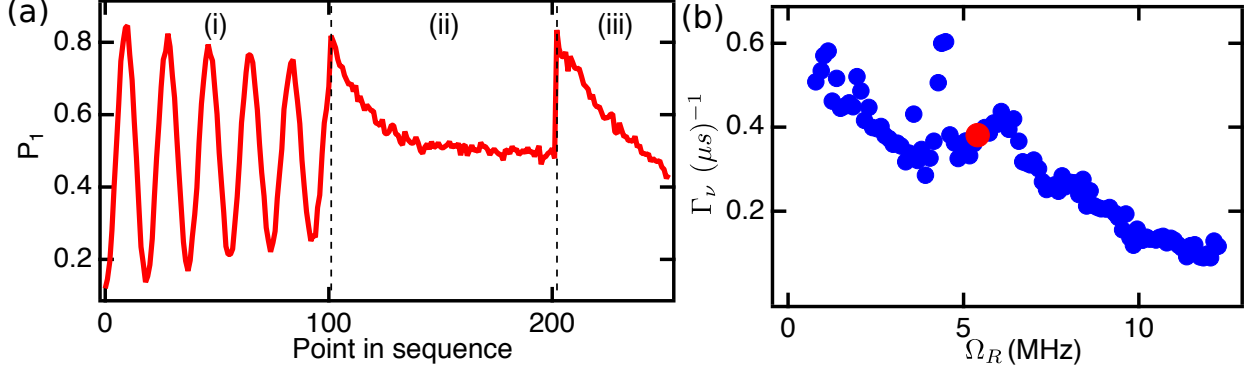


Figure B.3: **Spin-locking spectroscopy:** (a) A typical spin-locking measurement run, where we measure (i) a Rabi flop to extract Ω_R (data span $1\ \mu\text{s}$), (ii) spin-locking decay to extract $\Gamma_{1\rho}$ (data span $15\ \mu\text{s}$), and (iii) T_1 population decay to extract Γ_1 (data span $45\ \mu\text{s}$). We fit these data to extract the desired parameters, and calculate $\Gamma_\nu = \Gamma_{1\rho} - \Gamma_1/2$. (b) We repeat this measurement for many Rabi-driving amplitudes to build up a dataset of Γ_ν as a function of Ω_R (the measurement shown in (a) corresponds to the red point in (b).) In this particular measurement, we see a $1/f$ -like dependence, with several noise peaks near 5 MHz.

measure Γ_1 in parallel so that we can subtract it off of $\Gamma_{1\rho}$. We also need to know what the Rabi frequency of the applied driving tone is, since Ω_R determines the frequency of the noise spectral density we sample in a measurement. A typical measurement strategy is to write one long sequence that does all three measurements (inversion recovery, Rabi drive, and spin-locking) back-to-back-to-back, and interlacing these measurements before averaging. This strategy is useful since a typical spin-locking measurement consists of a loop over many spin-locking drive amplitudes, and T_1 may fluctuate significantly over the typical time required to take such a measurement ($\sim$ several hours.) An example of such a measurement is shown in Fig. B.3(a), which produces one point in a collection of Γ_ν 's plotted as a function of Ω_R in Fig. B.3(b).

As stated previously, a typical spin-locking measurement consists of a loop over spin-locking drive amplitudes. In order to execute such a loop, we must apply both fixed amplitude rotations around the Bloch sphere and a variable amplitude spin-locking drive. This

complicates the control pulse programming, since a two axis sweep (one axis varying the spin-locking time τ and one varying the drive amplitude) would overload the per-channel memory of the Tabor. To circumvent this problem, we use all four output channels on the Tabor to program control pulses: two channels (Fig. B.2(a), red) output $I(t)$ and $Q(t)$ for the fixed Bloch sphere rotations, and two (Fig. B.2(a), green) output $I(t)$ and $Q(t)$ for the spin-locking drive. We may then modify the amplitude of only the spin-locking channels without hardwiring the amplitude changes in the control pulse sequence. Since we use all four channels to control the qubit, we perform the measurement with a microwave switch, which is actuated with a marker channel on the Tabor (Fig. B.2(a), blue).

APPENDIX C

Fabrication Recipes

Substrate cleaning

SAW devices at $\sim$ GHz range frequencies tend to have large ($\sim$ 100's of μm to mm) areas covered with small ($\sim$ 100 nm) patterns. Fabricating these structures requires substrates to be *very* clean. A recipe I've found for cleaning substrates that generally works for preparing SAW devices:

1. Soak in remover PG on the hotplate for 10 minutes at 80°C .
2. Transfer to petri-dish full of acetone using *metal tipped tweezers*. PG will dissolve most carbon-tipped tweezer. Transfer again to dilute the PG if you're paranoid.
3. While in acetone, rub substrate with Ruby stick. Remove from acetone, and rinse with IPA before acetone evaporates from surface. Blow dry.
4. Sonicate in acetone for 5-15 min. Use the good sonicator in the Class 1000 fume hood. You can heat up the sonicator if you want, but I somewhat doubt the acetone gets appreciably warmer if you do.
5. Sonicate in IPA for 5-15 min, again using the good sonicator in the Class 1000 fume hood.

6. Remove sample from IPA and blow dry. I've noticed that if, rather than being blown off by the nitrogen gun, the IPA evaporates off, it will leave a residue that can mess up fab further down the line, so be diligent when doing this step.
7. Plasma clean in O₂ plasma etcher. 300 W for 300 s, sitting in a petri dish covered in Al foil (new foil if you're paranoid). The substrate will be hot to the touch after the etch, so it's best to bring it to the class 100 room fume hood and let it sit for several minutes before you deposit resist.

It's also important to remove organic residues from qubit devices, which can serve as sources of decoherence. We have not done a systematic study to see how effective this cleaning method is for preparing qubit devices, though that would be a good thing to do at some point.

Positive photolithography

Photolithography is the process of defining patterns on a substrate for metallization/etching by coating the substrate with a light-sensitive polymer and exposing certain areas to UV light. Depending on the polymer and process, the polymer layer in the exposed area is either removed (positive process) or retained while the unexposed polymer is removed (negative process). We outline a standard recipe for the positive photolithography process in our lab.

1. Clean the substrate, using § C as a guide.
2. Spincoat S1813 G2 photoresist onto the substrate, 5000 RPM, 50 seconds.
3. Bake on resist hotplate, 60 seconds 110 °C

4. Align photomask relative to substrate. Once aligned, bring the substrate in contact with the photomask. If you are exposing a large area, you should be able to see Fresnel rings form between the substrate and the bottom of the photomask.
5. Once satisfied with alignment and contact, expose to UV for 8 seconds
6. Remove substrate from mask aligner and soak in chlorobenzene for ~ 5 minutes. This hardens the upper layer of photoresist, giving an undercut with a single layer process and making it easier to lift off metal.
7. Develop in 352 developer for 10-25 seconds, DI for 10 seconds, and blow dry. Inspect under the optical microscope. Depending on how precise you need your features to be, if the exposed parts aren't fully developed you can stick the substrate back into 352 for 5-10 seconds (with DI and blow dry afterwards) to develop a little more. But don't tell anyone I told you that.

E-beam lithography

In the same way that the resolution of an optical microscope is ultimately limited by the diffraction of light passing through a small aperture, the feature size possible to define by photolithography is ultimately limited by diffraction. We expect this limit to be of the same order as the wavelength of the light in question: in photolithography, we use UV light with a wavelength of $\lambda \approx 400$ nm. While industrial processes have produced techniques that make sub-wavelength features possible, for our tabletop mask-aligner we're limited in practice to features with linear dimensions $> 2 - 3 \mu\text{m}$. Therefore, if we want to make smaller features, we must expose the resist with a smaller wavelength particle. The electrons emitted from a

scanning electron microscope (SEM) perform this function for us: we use a modified SEM to point and shoots the emitted electron beam at the areas of substrate we desire to expose, “writing” a pattern in the resist that is then developed away to expose substrate. Electron beam (e-beam) lithography can create much smaller features than photolithography, at the expense of a much longer write time, since only a small area can be exposed at any given time.

General bi-layer resist recipe

This a general, bi-layer resist e-beam lithography recipe which is a standard workhorse recipe in our lab. We use a bi-layer stack of resist with MMA on the bottom and PMMA on the top. MMA is a “softer” resist than PMMA in that it generally requires a lower dose to expose, which results in resist undercuts that make lift-off easier. We can also use this bi-layer to make resist bridges, which are crucial in double-angle evaporation techniques for making Josephson junctions.

1. If applicable, photolithographically define alignment marks. Marks must be thick enough to be seen through the resist: 50 nm thick gold alignment marks generally work.¹
2. Spin-coat MMA/EL-9 resist on a clean substrate, 4000 RPM for 45 seconds. If you want to write many patterns on a single chip, it is advisable to inspect the back of the substrate after spin-coating, since resist may pool on the back, which will result in the substrate being tilted during writing and may cause issues with focusing.

¹A “rule of thumb” I’ve found is that the alignment marks should have thickness $\geq 2600/(Z - 14)$, where Z is the atomic number of the metal the mark is made out of. I really don’t have any idea how accurate this rule is, but it produces numbers that have anecdotally worked for me: gold has $Z = 79$ and thus you want alignment marks ≥ 40 nm thick, copper has $Z = 29$ and so to make copper alignment marks you need ≥ 170 nm, while aluminum has $Z = 13$, and thus you can’t reliably make alignment marks out of aluminum.

3. Bake substrate on hotplate, 180° C for 10 minutes, or in the big oven at 180° C for 1 hour. Remove from hotplate and allow to cool
4. Spin-coat PMMA C2 resist, 4000 RPM for 45 seconds. Again advisable to inspect the back of the substrate.
5. Bake substrate on hotplate, 180° C for 10 minutes, or in the big oven at 180° C for 1 hour. Remove from hotplate and allow to cool
6. Secure the device to the SEM stage. Prior to inserting into the SEM/EBL system, make a small scratch on the corner of the sample to focus on.
7. Insert SEM stage and sample into the vacuum chamber and pump down. Once high vacuum is obtained, on the NPGS computer set the system to “SEM/imaging” mode.
8. Move to a position away from your substrate, and find a feature (dust/scratch/etc...) on the surface of the sample stage. Adjust focus/astigmatism until you get a clear image of the feature at 100,000× magnification.
9. Move to the scratch on the corner of the substrate, and refocus. At this point the beam should be ready for writing.
10. Switch to “Full EBL” mode and move to desired writing area. If needed, do alignment. Switching to “Full EBL” prior to moving to the write area will prevent unnecessary exposure. If you need to write a specific part of the substrate/align, I’ve found it convenient to make large, very visible alignment marks

Superconducting qubit fabrication

Fabrication of superconducting qubit devices in our lab was started not too long before this thesis was written, and as our local capabilities/knowledge change, the fabrication process has also been changing. In light of this, this section should be thought of as a snapshot of an evolving fabrication process. I will detail to the greatest extent possible the current process, however the reader should be deterred from believing anything here is a definitive (or even optimal) process.

Substrate preparation

The current iteration of the fabrication process starts off with whole 2” silicon wafers. The wafers are cleaned using steps 1-5 in the substrate cleaning recipe above (the Ruby stick cleaning may be omitted.) After sonication in IPA, DI water straight from the tap in the fume hood is run over the surface of the substrate for $\sim 30 - 60$ seconds, and then the wafer is blow dried.

We then spincoat MMA/PMMA onto the entire wafer, as detailed in steps 2-5 of the “General bi-layer resist recipe” section. Once this done we score the wafer with a scoring pen, and cleave it into $\sim \frac{1}{2}'' \times \frac{1}{2}''$ pieces for electron beam lithography. This process will introduce some detritus onto the surface, however the junction area, which is the only place that requires high precision fabrication, is small and the probability of a piece of silicon dust hitting the exact spot a junction is written is rather low.

Dolan bridge junction fabrication process

As stated in Chapter 3 of the main text, the Josephson junction is defined with electron beam lithography and a two-step evaporation process called the “Dolan bridge” process [80].

In the Dolan bridge process, a bi-layer resist stack consisting of a “hard” resist (PMMA/C2) sitting on a “soft” resist (MMA/EL9) is used. As shown in Fig. C.1(a), a high dose (blue, $\sim 300 \mu\text{C}/\text{cm}^2$) is used to expose the hard resist, while a low dose (orange, $\sim 50 \mu\text{C}/\text{cm}^2$) exposes the soft resist while leaving the hard resist intact. The resist is then developed in MIBK:IPA 1:3 for 50 seconds, IPA for 15 seconds, and blow dried. Upon development, in the places where a low dose was used, the hard resist will remain intact. Near the edges of the device, this “overhang” is useful for facilitating liftoff, however the important part is the center, where the MMA layer is cleared out from under a “bridge” of PMMA that will define the junction. Fig. C.1(b) details these features.

The substrate is then plasma etched in 500 mTorr O_2 at 100 W for 20 seconds to remove residual organics on the exposed surface. After etching, we place the substrate into a thermal evaporator with the ability to controllably tilt the sample holder, and pump down to low pressure.² Once the pressure is low enough, the sample holder is tilted such that the substrate makes an angle θ_1 relative to the evaporation source, with the rotation axes running parallel to the bridge (tilt axis perpendicular to the bridge), and a film of aluminum is evaporated, see Fig. C.1(c). After the first evaporation, the device is exposed to a 90/10 Ar/ O_2 environment, which oxidized the surface of the aluminum, forming the insulating barrier that will provide the tunnel barrier for the junction (dark grey in Fig. C.1(d).) The evaporator is then once again evacuated to low pressure, and a second layer of aluminum is deposited at angle θ_2 . As shown in Fig. C.1(e), if done correctly the two large metallic pads on either side of the bridge will be connected by only a small area oxide barrier directly under the bridge, highlighted in yellow. After evaporation, the substrate is placed in acetone overnight, which dissolves the resist and leaves behind only the pattern of metal on the substrate, see Fig. C.1(f). We then

²Our thermal evaporator, if operated properly, can reach $\sim 10^{-7}$ Torr in ~ 3 hours.

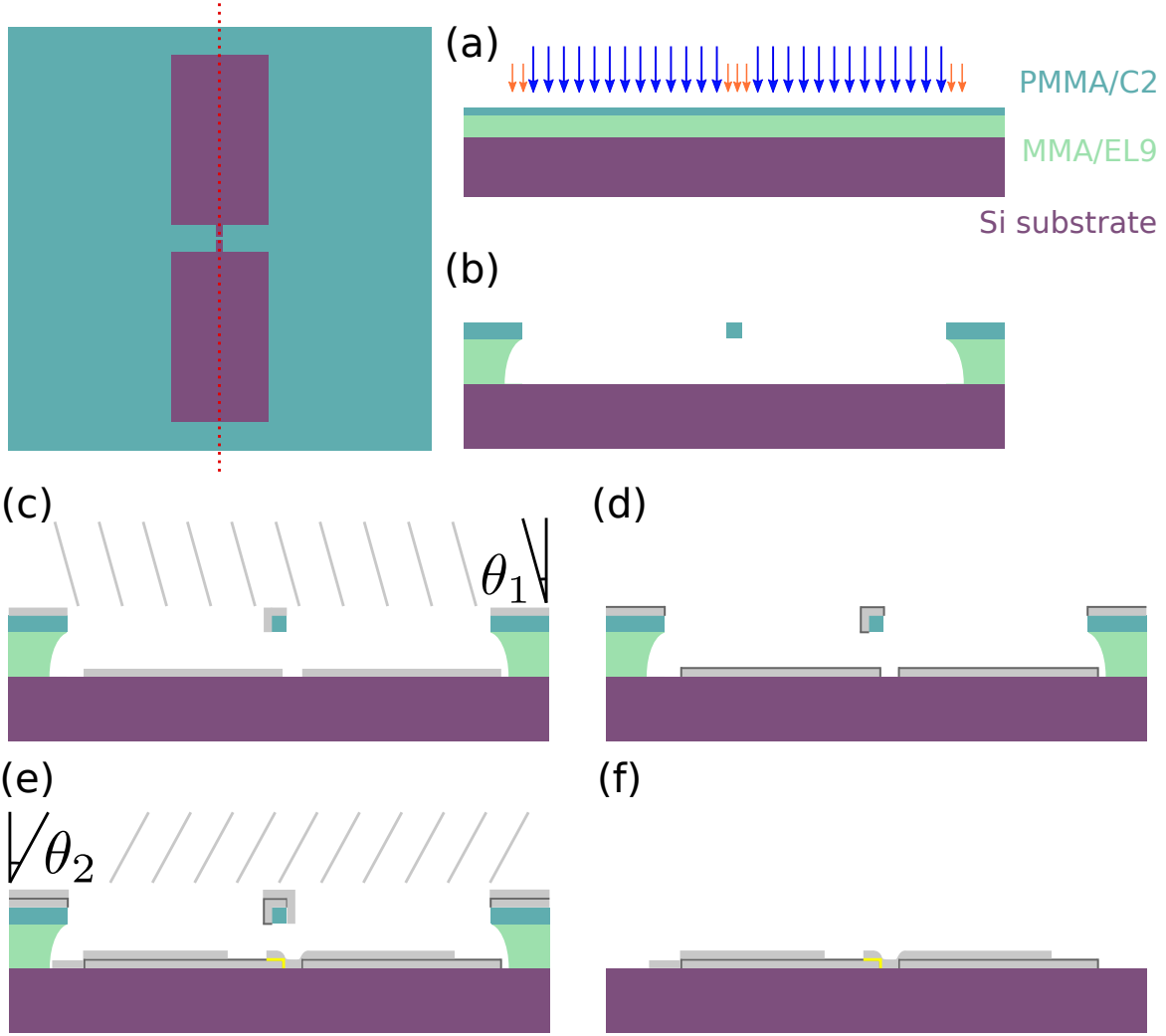


Figure C.1: **Description of the Dolan bridge junction fabrication process:** Process flow of the Dolan bridge junction fabrication process. Top-right: a top down view of the substrate with a resist bridge defined via electron beam lithography. All cross sections (a-f) run along the red dotted line. In panels (e-f), the part of the oxide layer that defines the junction is highlighted in yellow. Note that distances are not to scale.

clean the sample by briefly sonicating in acetone/IPA, and then blow drying and inspecting under the optical microscope.

The Josephson energy E_J of the junction is dictated by both the thickness of the oxide layer and the size of the overlap. The thickness of the layer will depend on how long and at what pressure the first metal layer is exposed to O_2 : we oxidize for 10 minutes at 3.4

Torr 90/10 Ar/O₂. Higher partial O₂ pressures, or longer exposure times, will increase the oxide layer thickness, decreasing the critical current and thus E_J . Since the tunneling rate will depend exponentially on the barrier thickness, this dependence will likely be nonlinear. Thus, we haven't done much optimization on this front, opting to stick with a recipe we know works.

Using the oxidation parameters described above, and a 100 nm×100 nm junction area, we fairly reliably fabricate junctions with $E_J/h \sim 8 - 10$ GHz. The critical current I_C , and thus E_J , is proportional to the junction area, so adjustments in the desired E_J can straightforwardly be made by modifying the bridge dimensions. The main quantity of interest in calculating the Junction area for a given geometry/set of angle θ_1 and θ_2 is the height of the “soft” resist layer. For MMA/EL9 spin-coated at 4000 RPM, the nominal thickness is 330 nm, though we have taken no real attempt to characterize this. As seen in Fig. C.1(e), the second metallization layer must be deposited on both the area above the first metallization and directly on the substrate. To ensure galvanic contact between metal deposited on these areas, it helps to make the second deposition layer appreciably thicker than the first: we generally deposit 20 nm of aluminum during the first evaporation and 60 nm during the second.

High-frequency SAW device fabrication

Surface acoustic wave devices with resonant frequencies in the 3 – 5 GHz range can be somewhat difficult to fabricate, as the wavelength, and thus the finger spacing, scales inversely with frequency. Most common piezoelectric substrates have a SAW propagation velocity $v_s \approx 3000$ m/s, so the wavelength of a 5 GHz SAW is $\lambda = f/v_s \approx 600$ nm. Since there are two IDT fingers per wavelength (one attached to each electrode), to fabricate a SAW device

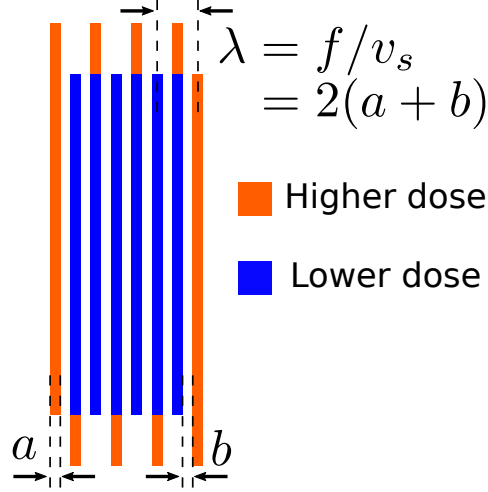


Figure C.2: **SAW device fabrication:** The principle wavelength of SAWs that an IDT can launch is defined by fabrication, $\lambda = 2(a + b)$ where a is the width of the IDT fingers and b is the spacing between fingers. When defining high-frequency SAW IDTs or Bragg mirrors, proximity dosing from neighboring fingers becomes significant. To compensate for this, it is helpful to lower the electron beam dose in for the IDTs away from the edges, and up the dose for the IDTs at the edge and for the sparsely spaced fingers at the edge of an IDT.

with equal electrode (a) and substrate (b) length we need to fabricate strips of 150 nm metal spaced by 150 nm ($a = b = 150$ nm) that are many 10's of microns long.

Clearly, in order to make these devices we have to use EBL. However, a bi-layer recipe is unavailable to us: a 100 – 200 nm wide strip of bi-layer MMA/PMMA that is many microns long will almost invariably collapse, since proximity dosing will cause an undercut of order the width the strip. Thus, in order to make high-frequency SAW devices, we are forced to use a single layer process.

Below is an outline of a process for reproducibly making SAW devices with $\lambda \geq 600$ nm using a positive pattern and metallization/liftoff procedure:

1. Clean off the substrate using the procedure outlined in section C. It is important to be rigorous about the cleaning for these devices, as Bragg mirrors and IDTs tend to be large structures, and one piece of debris in them can ruin an entire device.

2. Spincoat *only* PMMA following the recipe outlined in section C.
3. If using LiNbO₃ or ST-X quartz: deposit a ~ 30 nm aluminum discharging layer on top of the PMMA. The exact thickness doesn't seem to matter all that much: on quartz I've found that 10 nm is too thin, and I've accidentally deposited 60 nm and still was able to create functional devices. If fabricating on GaAs, this step is not necessary.
4. Following the instruction in section C, insert the substrate into the SEM. Take extra care focusing the beam, since the minimum feature sizes are small.
5. When designing the NPGS run-file, it is important to take proximity dosing into account. Electrons that hits the resist/discharging layer can scatter (forward scattering) and electrons that collide with the substrate can scatter back up to the resist layer (back scattering.) Either effect causes a finite dose away from the nominal writing area. Since at high frequencies SAW fingers are close to each other, proximity dosing plays a major roll. It is useful to raise the dose of leads near the edge of the sample, and in areas where the leads from only one electrode are (orange areas in Fig. C.2.)
6. In the run file, it is beneficial to make the nominal transducer width in the run-file a smaller than the desired width, and vary the dose until the desired width is produced. For example, if the desired $\lambda_{SAW} = 600$ nm, and the desired metallization ratio is $a = b = 150$ nm (see Fig. C.2), a good place to start is a run-file with $a = 100$ nm and $b = 200$ nm. Figure C.3 shows several examples of what typically happens when varying the dose.
7. Good starting doses: For GaAs, $270 \mu\text{C}/\text{cm}^2$ for inner areas (blue in Fig. C.2) and $320 \mu\text{C}/\text{cm}^2$ for outer areas (orange in Fig. C.2.) For ST-X quartz, $290 \mu\text{C}/\text{cm}^2$ for

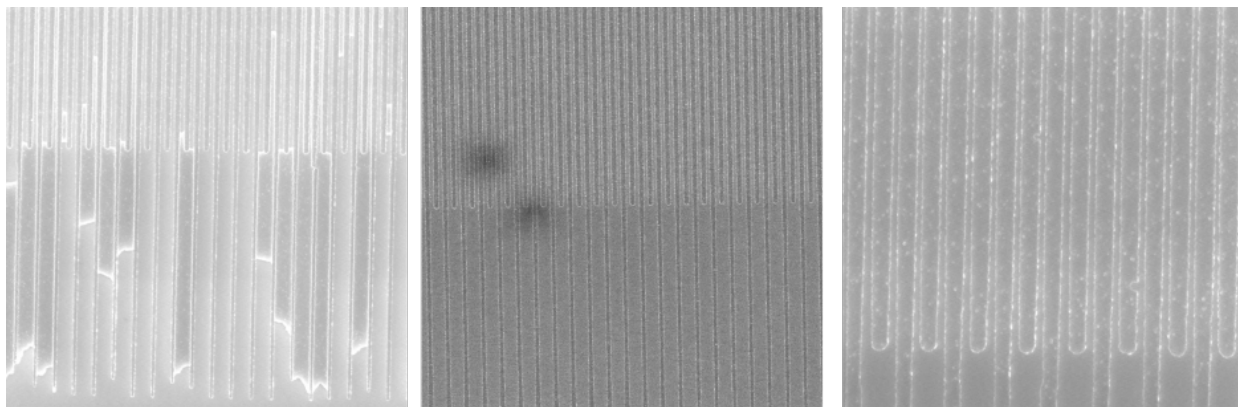


Figure C.3: **Examples of SAW IDT dosing effects:** Left: underdosed SAW IDT. Generally, when the IDTs are underdosed, lift-off is difficult. Center: optimized dose. Lift-off is easy, and $a \approx b$. Right: overdosed IDT. Clearly in this case, $a \gg b$, however even though the space between the IDT fingers is very small, lift-off generally goes well. Eventually, at too high a dose, lift off will get spotty right before the IDTs simply short together.

inner areas and $340 \mu\text{C}/\text{cm}^2$ for outer areas. These are suggestions that have worked for me in the past, not hard numbers. Dose will change based on geometry, resist age, magnification, etc...

8. Once writing is done, remove from EBL. If an aluminum discharging layer was used, remove the aluminum in a KOH developer such as AZ 300 MIF, rinse in DI and blow dry. Develop in 3:1 IPA:MIBK for 50 seconds, IPA for 15 seconds, and blow dry. Inspect under optical microscope.
9. Metallize. I have found that 30 nm of aluminum can reliably be lifted off, but thicker layers may be possible.
10. After evaporation, place the sample in a beaker of acetone and lift-off over night.
11. Try to remove as much visible metal as possible by pulling it off with tweezers/blowing the substrate with a stream of acetone (either using a syringe or just from the spray bottle.)

12. Fish out as much metal as possible from the acetone, or pour out most of the acetone and replenish it several times to get rid as much metal as possible in the beaker. It is important to **always** keep the surface of the substrate immersed under acetone until you're done with lift-off.
13. Sonicate for 5-10 minutes.
14. Sonication will break up the remaining metal into small particulates. Dilute metal particulates by again pouring out most of the acetone and replenishing it several times, again being careful not to expose the substrate to air.
15. While the substrate is still immersed in acetone, rub the area(s) of the substrate that have SAW devices using a ruby stick³. Do the rubbing motion along the IDT fingers. Rub for ≈ 30 seconds.
16. Remove the substrate from the beaker while spraying it with acetone. Before the acetone on the surface evaporates, rinse the substrate with IPA, and then blow dry.

Hard baked substrate spacers

Below is an outline of the process for making the substrate spacers used in the quantum acoustics flip-chip experiment. This recipe works well for 4 layers of photoresist, which renders bump bonds $\approx 4 \mu\text{m}$ thick.

1. Spincoat S1813 G2 resist, 5000 RPM for 50 seconds
2. Bake substrate at 110°C for 1 minute

³If you decide you want to use remover PG to do lift-off, make sure you switch to acetone by this step, since PG dissolve the tips of Ruby sticks.

3. Repeat steps (1) and (2) 3 more times for a total of 4 layers.
4. Expose for 90 seconds
5. Develop in 352 developer for ~ 30 -60 seconds. Note that it might be a good idea to use a larger beaker to dilute the photoresist, since 4 layers is a lot of resist. Rinse in DI for 15 seconds, and then blow dry. Remember that 352 is a NaOH based developer, and that NaOH also etches aluminum, so it is better to increase the exposure time if the development is slow.
6. Inspect under the optical microscope. I have found that the last exposed resist to come off is on the pads of the device, so make sure resist is stripped clean from there.
7. Optional: Oxygen plasma etch at 300 W to remove any residual resist. Plasma etching will etch the spacers as well, at a rate of ≈ 200 nm/s, so be careful to not over-etch. Four layers of S1813 G2 nominally gives you $4.6\text{ }\mu\text{m}$ anyways, so if you're aiming for as close to $4\text{ }\mu\text{m}$ as possible, etch for 180 seconds.
8. Put on hotplate, which is still at $110\text{ }^{\circ}\text{C}$. Set the hotplate to $250\text{ }^{\circ}\text{C}$, and let bake for 1.5 hours.

Note that, when spin-coating on small substrates, resist can pool to significant depth on the edge of the substrate, which will cause unreliable spacing. My strategy to avoid this resist is twofold: fabricate multiple SAW devices on a larger substrate, and also leave some sacrificial area near the edge of the substrate designed to be diced off after spacer fabrication. For example, I generally fabricate 6 devices on one substrate: this would require a 14×21 mm substrate, however I actually use a 16×22 mm substrate and cleave the edges off when dicing the larger substrate to individual device size.

BIBLIOGRAPHY

BIBLIOGRAPHY

- [1] P. W. Anderson. More is different. *Science*, 177(4047):393–396, 1972.
- [2] A. D. O’Connell, M. Hofheinz, M. Ansmann, Radoslaw C. Bialczak, M. Lenander, Erik Lucero, M. Neeley, D. Sank, H. Wang, M. Weides, J. Wenner, John M. Martinis, and A. N. Cleland. Quantum ground state and single-phonon control of a mechanical resonator. *Nature*, 464:697 EP –, 03 2010.
- [3] J. M. Pirkkalainen, S. U. Cho, Jian Li, G. S. Paraoanu, P. J. Hakonen, and M. A. Sillanpää. Hybrid circuit cavity quantum electrodynamics with a micromechanical resonator. *Nature*, 494:211 EP –, 02 2013.
- [4] Patricio Arrangoiz-Arriola, E. Alex Wollack, Zhaoyou Wang, Marek Pechal, Wentao Jiang, Timothy P. McKenna, Jeremy D. Witmer, Raphaël Van Laer, and Amir H. Safavi-Naeini. Resolving the energy levels of a nanomechanical oscillator. *Nature*, 571(7766):537–540, 2019.
- [5] Yiwen Chu, Prashanta Kharel, William H. Renninger, Luke D. Burkhardt, Luigi Frunzio, Peter T. Rakich, and Robert J. Schoelkopf. Quantum acoustics with superconducting qubits. *Science*, 2017.
- [6] Mikael Kervinen, Ilkka Rissanen, and Mika Sillanpää. Interfacing planar superconducting qubits with high overtone bulk acoustic phonons. *Phys. Rev. B*, 97:205443, May 2018.
- [7] Martin V. Gustafsson, Thomas Aref, Anton Frisk Kockum, Maria K. Ekström, Göran Johansson, and Per Delsing. Propagating phonons coupled to an artificial atom. *Science*, 346(6206):207–211, 2014.
- [8] Riccardo Manenti, Anton F. Kockum, Andrew Patterson, Tanja Behrle, Joseph Rahamim, Giovanna Tancredi, Franco Nori, and Peter J. Leek. Circuit quantum acoustodynamics with surface acoustic waves. *Nature Communications*, 8(1):975, 2017.
- [9] Atsushi Noguchi, Rekishu Yamazaki, Yutaka Tabuchi, and Yasunobu Nakamura. Qubit-assisted transduction for a detection of surface acoustic waves near the quantum limit. *Phys. Rev. Lett.*, 119:180505, Nov 2017.
- [10] K. J. Satzinger, Y. P. Zhong, H. S. Chang, G. A. Peairs, A. Bienfait, Ming-Han Chou, A. Y. Cleland, C. R. Conner, É. Dumur, J. Grebel, I. Gutierrez, B. H. November, R. G. Povey, S. J. Whiteley, D. D. Awschalom, D. I. Schuster, and A. N. Cleland. Quantum control of surface acoustic-wave phonons. *Nature*, 563(7733):661–665, 2018.

- [11] Bradley A. Moores, Lucas R. Sletten, Jeremie J. Viennot, and K. W. Lehnert. Cavity quantum acoustic device in the multimode strong coupling regime. *Phys. Rev. Lett.*, 120:227701, May 2018.
- [12] L. R. Sletten, B. A. Moores, J. J. Viennot, and K. W. Lehnert. Resolving phonon Fock states in a multimode cavity with a double-slit qubit. *Phys. Rev. X*, 9:021056, Jun 2019.
- [13] Yiwen Chu, Prashanta Kharel, Taekwan Yoon, Luigi Frunzio, Peter T. Rakich, and Robert J. Schoelkopf. Creation and control of multi-phonon Fock states in a bulk acoustic-wave resonator. *Nature*, 563(7733):666–670, 2018.
- [14] X. Ma, J. J. Viennot, S. Kotler, J. D. Teufel, and K. W. Lehnert. Non-classical energy squeezing of a macroscopic mechanical oscillator. *Nature Physics*, 2021.
- [15] A. Bienfait, Y. P. Zhong, H.-S. Chang, M.-H. Chou, C. R. Conner, É. Dumur, J. Grebel, G. A. Peairs, R. G. Povey, K. J. Satzinger, and A. N. Cleland. Quantum erasure using entangled surface acoustic phonons. *Phys. Rev. X*, 10:021055, Jun 2020.
- [16] Anton Frisk Kockum, Per Delsing, and Göran Johansson. Designing frequency-dependent relaxation rates and lamb shifts for a giant artificial atom. *Phys. Rev. A*, 90:013837, Jul 2014.
- [17] Gustav Andersson, Baladitya Suri, Lingzhen Guo, Thomas Aref, and Per Delsing. Non-exponential decay of a giant artificial atom. *Nature Physics*, 15(11):1123–1127, 2019.
- [18] Connor T. Hann, Chang-Ling Zou, Yaxing Zhang, Yiwen Chu, Robert J. Schoelkopf, S. M. Girvin, and Liang Jiang. Hardware-efficient quantum random access memory with hybrid quantum acoustic systems. *Phys. Rev. Lett.*, 123:250501, Dec 2019.
- [19] Mohammad Mirhosseini, Alp Sipahigil, Mahmoud Kalaei, and Oskar Painter. Superconducting qubit to optical photon transduction. *Nature*, 588(7839):599–603, 2020.
- [20] H. K. Onnes. *Leiden. Comm.*, 120b, 122b, 124c, 1911.
- [21] M. Tinkham. *Introduction to Superconductivity*. Dover Books on Physics Series. Dover Publications, 2004.
- [22] V. L. Ginzberg and L. D. Landau. On the theory of superconductivity. *Soviet Physics JETP*, 20:1064, 1950.
- [23] J. Bardeen, L. N. Cooper, and J. R. Schrieffer. Theory of superconductivity. *Phys. Rev.*, 108:1175–1204, Dec 1957.

- [24] L. P. Gor'kov. Microscopic derivation of the Ginzberg-Landau equations in the theory of superconductivity. *Soviet Physics JETP*, 36(9):1918–1923, June 1959.
- [25] B.D. Josephson. Possible new effects in superconductive tunnelling. *Physics Letters*, 1(7):251 – 253, 1962.
- [26] P. W. Anderson and J. M. Rowell. Probable observation of the Josephson superconducting tunneling effect. *Phys. Rev. Lett.*, 10:230–232, Mar 1963.
- [27] A. J. Leggett. Macroscopic Quantum Systems and the Quantum Theory of Measurement. *Progress of Theoretical Physics Supplement*, 69:80–100, 03 1980.
- [28] John Clarke, Andrew N. Cleland, Michel H. Devoret, Daniel Esteve, and John M. Martinis. Quantum mechanics of a macroscopic variable: The phase difference of a Josephson junction. *Science*, 239(4843):992–997, 1988.
- [29] Y. Nakamura, Yu. A. Pashkin, and J. S. Tsai. Coherent control of macroscopic quantum states in a single-Cooper-pair box. *Nature*, 398(6730):786–788, 1999.
- [30] N. K. Langford. Circuit QED - lecture notes. arXiv:1310.1897v1 [quant-ph], 10 2013.
- [31] M. H. Devoret, B Huard, R. J. Schoelkopf, and L. F. Cugliandolo. *Quantum Machines: Measurement and Control of Engineered Quantum Systems: Lecture Notes of the Les Houches Summer School*, volume 96. Oxford University Press, 2014.
- [32] Uri Vool and Michel Devoret. Introduction to quantum electromagnetic circuits. *International Journal of Circuit Theory and Applications*, 45(7):897–934, June 2017.
- [33] R. J. Schoelkopf and S. M. Girvin. Wiring up quantum systems. *Nature*, 451:664 EP –, 02 2008.
- [34] Jens Koch, Terri M. Yu, Jay Gambetta, A. A. Houck, D. I. Schuster, J. Majer, Alexandre Blais, M. H. Devoret, S. M. Girvin, and R. J. Schoelkopf. Charge-insensitive qubit design derived from the cooper pair box. *Phys. Rev. A*, 76:042319, Oct 2007.
- [35] R. J. Schoelkopf, P. Wahlgren, A. A. Kozhevnikov, P. Delsing, and D. E. Prober. The radio-frequency single-electron transistor (RF-SET): A fast and ultrasensitive electrometer. *Science*, 280(5367):1238–1242, 1998.
- [36] D.J. Griffiths and P.D.J. Griffiths. *Introduction to Quantum Mechanics*. Pearson international edition. Pearson Prentice Hall, 2005.
- [37] D. I. Schuster, A. A. Houck, J. A. Schreier, A. Wallraff, J. M. Gambetta, A. Blais, L. Frunzio, J. Majer, B. Johnson, M. H. Devoret, S. M. Girvin, and R. J. Schoelkopf.

- Resolving photon number states in a superconducting circuit. *Nature*, 445:515 EP –, 02 2007.
- [38] A. A. Houck, D. I. Schuster, J. M. Gambetta, J. A. Schreier, B. R. Johnson, J. M. Chow, L. Frunzio, J. Majer, M. H. Devoret, S. M. Girvin, and R. J. Schoelkopf. Generating single microwave photons in a circuit. *Nature*, 449:328 EP –, 09 2007.
 - [39] J. Majer, J. M. Chow, J. M. Gambetta, Jens Koch, B. R. Johnson, J. A. Schreier, L. Frunzio, D. I. Schuster, A. A. Houck, A. Wallraff, A. Blais, M. H. Devoret, S. M. Girvin, and R. J. Schoelkopf. Coupling superconducting qubits via a cavity bus. *Nature*, 449:443 EP –, 09 2007.
 - [40] Hanhee Paik, D. I. Schuster, Lev S. Bishop, G. Kirchmair, G. Catelani, A. P. Sears, B. R. Johnson, M. J. Reagor, L. Frunzio, L. I. Glazman, S. M. Girvin, M. H. Devoret, and R. J. Schoelkopf. Observation of high coherence in Josephson junction qubits measured in a three-dimensional circuit QED architecture. *Phys. Rev. Lett.*, 107:240501, Dec 2011.
 - [41] Chad Rigetti, Jay M. Gambetta, Stefano Poletto, B. L. T. Plourde, Jerry M. Chow, A. D. Córcoles, John A. Smolin, Seth T. Merkel, J. R. Rozen, George A. Keefe, Mary B. Rothwell, Mark B. Ketchen, and M. Steffen. Superconducting qubit in a waveguide cavity with a coherence time approaching 0.1 ms. *Phys. Rev. B*, 86:100506, Sep 2012.
 - [42] K. Serniak, S. Diamond, M. Hays, V. Fatemi, S. Shankar, L. Frunzio, R.J. Schoelkopf, and M.H. Devoret. Direct dispersive monitoring of charge parity in offset-charge-sensitive transmons. *Phys. Rev. Applied*, 12:014052, Jul 2019.
 - [43] M. S. Blok, V. V. Ramasesh, T. Schuster, K. O’Brien, J. M. Kreikebaum, D. Dahlen, A. Morvan, B. Yoshida, N. Y. Yao, and I. Siddiqi. Quantum information scrambling on a superconducting qutrit processor. *Phys. Rev. X*, 11:021010, Apr 2021.
 - [44] Alexander P. M. Place, Lila V. H. Rodgers, Pranav Mundada, Basil M. Smitham, Mattias Fitzpatrick, Zhaoqi Leng, Anjali Premkumar, Jacob Bryon, Andrei Vrajitoarea, Sara Sussman, Guangming Cheng, Trisha Madhavan, Harshvardhan K. Babla, Xuan Hoang Le, Youqi Gang, Berthold Jäck, András Gyenis, Nan Yao, Robert J. Cava, Nathalie P. de Leon, and Andrew A. Houck. New material platform for superconducting transmon qubits with coherence times exceeding 0.3 milliseconds. *Nature Communications*, 12(1):1779, 2021.
 - [45] J. M. Chow, J. M. Gambetta, L. Tornberg, Jens Koch, Lev S. Bishop, A. A. Houck, B. R. Johnson, L. Frunzio, S. M. Girvin, and R. J. Schoelkopf. Randomized benchmarking and process tomography for gate errors in a solid-state qubit. *Phys. Rev. Lett.*, 102:090502, Mar 2009.

- [46] L. DiCarlo, J. M. Chow, J. M. Gambetta, Lev S. Bishop, B. R. Johnson, D. I. Schuster, J. Majer, A. Blais, L. Frunzio, S. M. Girvin, and R. J. Schoelkopf. Demonstration of two-qubit algorithms with a superconducting quantum processor. *Nature*, 460(7252):240–244, 2009.
- [47] Jerry M. Chow, Jay M. Gambetta, Easwar Magesan, David W. Abraham, Andrew W. Cross, B R Johnson, Nicholas A. Masluk, Colm A. Ryan, John A. Smolin, Srikanth J. Srinivasan, and M Steffen. Implementing a strand of a scalable fault-tolerant quantum computing fabric. *Nature Communications*, 5:4015 EP –, 06 2014.
- [48] J. Kelly, R. Barends, A. G. Fowler, A. Megrant, E. Jeffrey, T. C. White, D. Sank, J. Y. Mutus, B. Campbell, Yu Chen, Z. Chen, B. Chiaro, A. Dunsworth, I. C. Hoi, C. Neill, P. J. J. O’Malley, C. Quintana, P. Roushan, A. Vainsencher, J. Wenner, A. N. Cleland, and John M. Martinis. State preservation by repetitive error detection in a superconducting quantum circuit. *Nature*, 519:66 EP –, 03 2015.
- [49] D. Ristè, S. Poletto, M. Z. Huang, A. Bruno, V. Vesterinen, O. P. Saira, and L. DiCarlo. Detecting bit-flip errors in a logical qubit using stabilizer measurements. *Nature Communications*, 6:6983 EP –, 04 2015.
- [50] Nissim Ofek, Andrei Petrenko, Reinier Heeres, Philip Reinhold, Zaki Leghtas, Brian Vlastakis, Yehan Liu, Luigi Frunzio, S. M. Girvin, L. Jiang, Mazyar Mirrahimi, M. H. Devoret, and R. J. Schoelkopf. Extending the lifetime of a quantum bit with error correction in superconducting circuits. *Nature*, 536:441 EP –, 07 2016.
- [51] P. J. J. O’Malley, R. Babbush, I. D. Kivlichan, J. Romero, J. R. McClean, R. Barends, J. Kelly, P. Roushan, A. Tranter, N. Ding, B. Campbell, Y. Chen, Z. Chen, B. Chiaro, A. Dunsworth, A. G. Fowler, E. Jeffrey, E. Lucero, A. Megrant, J. Y. Mutus, M. Neeley, C. Neill, C. Quintana, D. Sank, A. Vainsencher, J. Wenner, T. C. White, P. V. Coveney, P. J. Love, H. Neven, A. Aspuru-Guzik, and J. M. Martinis. Scalable quantum simulation of molecular energies. *Phys. Rev. X*, 6:031007, Jul 2016.
- [52] R. K. Naik, N. Leung, S. Chakram, Peter Groszkowski, Y. Lu, N. Earnest, D. C. McKay, Jens Koch, and D. I. Schuster. Random access quantum information processors using multimode circuit quantum electrodynamics. *Nature Communications*, 8(1):1904, 2017.
- [53] Abhinav Kandala, Antonio Mezzacapo, Kristan Temme, Maika Takita, Markus Brink, Jerry M. Chow, and Jay M. Gambetta. Hardware-efficient variational quantum eigensolver for small molecules and quantum magnets. *Nature*, 549(7671):242–246, 2017.
- [54] Frank Arute et al. Quantum supremacy using a programmable superconducting processor. *Nature*, 574(7779):505–510, 2019.

- [55] Mohammad Mirhosseini, Eunjong Kim, Xueyue Zhang, Alp Sipahigil, Paul B. Dieterle, Andrew J. Keller, Ana Asenjo-Garcia, Darrick E. Chang, and Oskar Painter. Cavity quantum electrodynamics with atom-like mirrors. *Nature*, 569(7758):692–697, 2019.
- [56] Brian Vlastakis, Gerhard Kirchmair, Zaki Leghtas, Simon E. Nigg, Luigi Frunzio, S. M. Girvin, Mazyar Mirrahimi, M. H. Devoret, and R. J. Schoelkopf. Deterministically encoding quantum information using 100-photon Schrödinger cat states. *Science*, 342(6158):607–610, 2013.
- [57] K. W. Murch, S. J. Weber, C. Macklin, and I. Siddiqi. Observing single quantum trajectories of a superconducting quantum bit. *Nature*, 502:211 EP –, 10 2013.
- [58] Z. K. Mineev, S. O. Mundhada, S. Shankar, P. Reinhold, R. Gutiérrez-Jáuregui, R. J. Schoelkopf, M. Mirrahimi, H. J. Carmichael, and M. H. Devoret. To catch and reverse a quantum jump mid-flight. *Nature*, 2019.
- [59] M. Naghiloo, J. J. Alonso, A. Romito, E. Lutz, and K. W. Murch. Information gain and loss for a quantum Maxwell’s demon. *Phys. Rev. Lett.*, 121:030604, Jul 2018.
- [60] Alexandre Blais, Ren-Shou Huang, Andreas Wallraff, S. M. Girvin, and R. J. Schoelkopf. Cavity quantum electrodynamics for superconducting electrical circuits: An architecture for quantum computation. *Phys. Rev. A*, 69:062320, Jun 2004.
- [61] A. Wallraff, D. I. Schuster, A. Blais, L. Frunzio, R. S. Huang, J. Majer, S. Kumar, S. M. Girvin, and R. J. Schoelkopf. Strong coupling of a single photon to a superconducting qubit using circuit quantum electrodynamics. *Nature*, 431:162 EP –, 09 2004.
- [62] D.P. DiVincenzo, Institut für Festkörperforschung, Institut für Festkörperforschung (Jülich). Spring School, and Jülich Centre for Neutron Science. *Quantum Information Processing: Lecture Notes of the 44th IFF Spring School 2013*. Schriften des Forschungszentrums Jülich. Forschungszentrum, Zentralbibliothek, 2013.
- [63] M. H. Devoret and R. J. Schoelkopf. Superconducting circuits for quantum information: An outlook. *Science*, 339(6124):1169–1174, 2013.
- [64] Alexandre Blais, Steven M. Girvin, and William D. Oliver. Quantum information processing and quantum optics with circuit quantum electrodynamics. *Nature Physics*, 16(3):247–256, 2020.
- [65] A. A. Houck, J. A. Schreier, B. R. Johnson, J. M. Chow, Jens Koch, J. M. Gambetta, D. I. Schuster, L. Frunzio, M. H. Devoret, S. M. Girvin, and R. J. Schoelkopf. Controlling the spontaneous emission of a superconducting transmon qubit. *Phys. Rev. Lett.*, 101:080502, Aug 2008.

- [66] P. Krantz, M. Kjaergaard, F. Yan, T. P. Orlando, S. Gustavsson, and W. D. Oliver. A quantum engineer's guide to superconducting qubits. *Applied Physics Reviews*, 6(2):021318, 2019.
- [67] Daniel Sank. *Fast, Accurate State Measurement in Superconducting Qubits*. PhD thesis, University of California Santa Barbara, 2014.
- [68] Jay Gambetta, Alexandre Blais, D. I. Schuster, A. Wallraff, L. Frunzio, J. Majer, M. H. Devoret, S. M. Girvin, and R. J. Schoelkopf. Qubit-photon interactions in a cavity: Measurement-induced dephasing and number splitting. *Phys. Rev. A*, 74:042318, Oct 2006.
- [69] David Schuster. *Circuit Quantum Electrodynamics*. PhD thesis, Yale, 2007.
- [70] Vladimir B. Braginsky, Farid Ya Khalili, and Kip S. Thorne. *Quantum Measurement*. Cambridge University Press, 1992.
- [71] Daniel Slichter. *Quantum Jumps and Measurement Backaction in a Superconducting Qubit*. PhD thesis, University of California Berkeley, 2011.
- [72] J. A. Schreier, A. A. Houck, Jens Koch, D. I. Schuster, B. R. Johnson, J. M. Chow, J. M. Gambetta, J. Majer, L. Frunzio, M. H. Devoret, S. M. Girvin, and R. J. Schoelkopf. Suppressing charge noise decoherence in superconducting charge qubits. *Phys. Rev. B*, 77:180502, May 2008.
- [73] Morten Kjaergaard, Mollie E. Schwartz, Jochen Braumüller, Philip Krantz, Joel I.-J. Wang, Simon Gustavsson, and William D. Oliver. Superconducting qubits: Current state of play. *Annual Review of Condensed Matter Physics*, 11(1):369–395, 2020.
- [74] David M Pozar. *Microwave engineering; 3rd ed.* Wiley, Hoboken, NJ, 2005.
- [75] D.F. Walls and G.J. Milburn. *Quantum Optics*. Springer Berlin Heidelberg, 2008.
- [76] A. P. Sears, A. Petrenko, G. Catelani, L. Sun, Hanhee Paik, G. Kirchmair, L. Frunzio, L. I. Glazman, S. M. Girvin, and R. J. Schoelkopf. Photon shot noise dephasing in the strong-dispersive limit of circuit QED. *Phys. Rev. B*, 86:180504, Nov 2012.
- [77] M. Göppl, A. Fragner, M. Baur, R. Bianchetti, S. Filipp, J. M. Fink, P. J. Leek, G. Puebla, L. Steffen, and A. Wallraff. Coplanar waveguide resonators for circuit quantum electrodynamics. *Journal of Applied Physics*, 104(11):113904, 2008.
- [78] Simon E. Nigg, Hanhee Paik, Brian Vlastakis, Gerhard Kirchmair, S. Shankar, Luigi Frunzio, M. H. Devoret, R. J. Schoelkopf, and S. M. Girvin. Black-box superconducting circuit quantization. *Phys. Rev. Lett.*, 108:240502, Jun 2012.

- [79] C. Wang, C. Axline, Y. Y. Gao, T. Brecht, Y. Chu, L. Frunzio, M. H. Devoret, and R. J. Schoelkopf. Surface participation and dielectric loss in superconducting qubits. *Applied Physics Letters*, 107(16):162601, 2015.
- [80] G. J. Dolan. Offset masks for lift-off photoprocessing. *Applied Physics Letters*, 31(5):337–339, 1977.
- [81] J M Kreikebaum, K P O’Brien, A Morvan, and I Siddiqi. Improving wafer-scale Josephson junction resistance variation in superconducting quantum coherent circuits. *Superconductor Science and Technology*, 33(6):06LT02, may 2020.
- [82] Jared B. Hertzberg, Eric J. Zhang, Sami Rosenblatt, Easwar Magesan, John A. Smolin, Jeng-Bang Yau, Vivekananda P. Adiga, Martin Sandberg, Markus Brink, Jerry M. Chow, and Jason S. Orcutt. Laser-annealing Josephson junctions for yielding scaled-up superconducting quantum processors. arXiv preprint, 2020.
- [83] Richard Phillips Feynman, Robert Benjamin Leighton, and Matthew Sands. *The Feynman lectures on physics; New millennium ed.* Basic Books, New York, NY, 2010.
- [84] J M Kreikebaum, A Dove, W Livingston, E Kim, and I Siddiqi. Optimization of infrared and magnetic shielding of superconducting TiN and al coplanar microwave resonators. *Superconductor Science and Technology*, 29(10):104002, aug 2016.
- [85] R. Barends, J. Wenner, M. Lenander, Y. Chen, R. C. Bialczak, J. Kelly, E. Lucero, P. O’Malley, M. Mariantoni, D. Sank, H. Wang, T. C. White, Y. Yin, J. Zhao, A. N. Cleland, John M. Martinis, and J. J. A. Baselmans. Minimizing quasiparticle generation from stray infrared light in superconducting quantum circuits. *Applied Physics Letters*, 99(11):113507, 2011.
- [86] M. J. Persky. Review of black surfaces for space-borne infrared systems. *Review of Scientific Instruments*, 70(5):2193–2217, 1999.
- [87] Jen-Hao Yeh, Jay LeFebvre, Shavindra Premaratne, F. C. Wellstood, and B. S. Palmer. Microwave attenuators for use with quantum devices below 100 mK. *Journal of Applied Physics*, 121(22):224501, 2017.
- [88] Fei Yan, Dan Campbell, Philip Krantz, Morten Kjaergaard, David Kim, Jonilyn L. Yoder, David Hover, Adam Sears, Andrew J. Kerman, Terry P. Orlando, Simon Gustavsson, and William D. Oliver. Distinguishing coherent and thermal photon noise in a circuit quantum electrodynamical system. *Phys. Rev. Lett.*, 120:260504, Jun 2018.
- [89] Frank Pobell. *Matter and Methods at Low Temperatures*. Springer-Verlag Berlin Heidelberg, 3 edition, 2007.

- [90] Z. Wang, S. Shankar, Z.K. Mineev, P. Campagne-Ibarcq, A. Narla, and M.H. Devoret. Cavity attenuators for superconducting qubits. *Phys. Rev. Applied*, 11:014031, Jan 2019.
- [91] S. Krinner, S. Storz, P. Kurpiers, P. Magnard, J. Heinsoo, R. Keller, J. Lütolf, C. Eichler, and A. Wallraff. Engineering cryogenic setups for 100-qubit scale superconducting circuit systems. *EPJ Quantum Technology*, 6(1):2, May 2019.
- [92] D F Santavicca and D E Prober. Impedance-matched low-pass stripline filters. *Measurement Science and Technology*, 19(8):087001, jun 2008.
- [93] A. A. Clerk, M. H. Devoret, S. M. Girvin, Florian Marquardt, and R. J. Schoelkopf. Introduction to quantum noise, measurement, and amplification. *Rev. Mod. Phys.*, 82:1155–1208, Apr 2010.
- [94] C. D. Motchenbacher and J. A. Connelly. *Low-Noise Electronic System Design*. John Wiley and Sons, Inc., 1993.
- [95] Lukas Grünhaupt, Nataliya Maleeva, Sebastian T. Skacel, Martino Calvo, Florence Levy-Bertrand, Alexey V. Ustinov, Hannes Rotzinger, Alessandro Monfardini, Gianluigi Catelani, and Ioan M. Pop. Loss mechanisms and quasiparticle dynamics in superconducting microwave resonators made of thin-film granular aluminum. *Phys. Rev. Lett.*, 121:117001, Sep 2018.
- [96] E. T. Mannila, P. Samuelsson, S. Simbierowicz, J. T. Peltonen, V. Vesterinen, L. Grönberg, J. Hassel, V. F. Maisi, and J. P. Pekola. A superconductor free of quasiparticles for seconds, 2021.
- [97] Lev S. Bishop, Eran Ginossar, and S. M. Girvin. Response of the strongly driven Jaynes-Cummings oscillator. *Phys. Rev. Lett.*, 105:100505, Sep 2010.
- [98] Maxime Boissonneault, J. M. Gambetta, and Alexandre Blais. Improved superconducting qubit readout by qubit-induced nonlinearities. *Phys. Rev. Lett.*, 105:100504, Sep 2010.
- [99] M. D. Reed, L. DiCarlo, B. R. Johnson, L. Sun, D. I. Schuster, L. Frunzio, and R. J. Schoelkopf. High-fidelity readout in circuit quantum electrodynamics using the Jaynes-Cummings nonlinearity. *Phys. Rev. Lett.*, 105:173601, Oct 2010.
- [100] D. I. Schuster, A. Wallraff, A. Blais, L. Frunzio, R.-S. Huang, J. Majer, S. M. Girvin, and R. J. Schoelkopf. ac Stark shift and dephasing of a superconducting qubit strongly coupled to a cavity field. *Phys. Rev. Lett.*, 94:123602, Mar 2005.
- [101] Alexandre Blais, Arne L. Grimsmo, S. M. Girvin, and Andreas Wallraff. Circuit quantum electrodynamics. *Rev. Mod. Phys.*, 93:025005, May 2021.

- [102] J.R. Johansson, P.D. Nation, and Franco Nori. Qutip 2: A python framework for the dynamics of open quantum systems. *Computer Physics Communications*, 184(4):1234 – 1240, 2013.
- [103] A. Wallraff, D. I. Schuster, A. Blais, L. Frunzio, J. Majer, M. H. Devoret, S. M. Girvin, and R. J. Schoelkopf. Approaching unit visibility for control of a superconducting qubit with dispersive readout. *Phys. Rev. Lett.*, 95:060501, Aug 2005.
- [104] G. Ithier, E. Collin, P. Joyez, P. J. Meeson, D. Vion, D. Esteve, F. Chiarello, A. Shnirman, Y. Makhlin, J. Schrieffer, and G. Schön. Decoherence in a superconducting quantum bit circuit. *Phys. Rev. B*, 72:134519, Oct 2005.
- [105] Łukasz Cywiński, Roman M. Lutchyn, Cody P. Nave, and S. Das Sarma. How to enhance dephasing time in superconducting qubits. *Phys. Rev. B*, 77:174509, May 2008.
- [106] Jonas Bylander, Simon Gustavsson, Fei Yan, Fumiki Yoshihara, Khalil Harrabi, George Fitch, David G. Cory, Yasunobu Nakamura, Jaw-Shen Tsai, and William D. Oliver. Noise spectroscopy through dynamical decoupling with a superconducting flux qubit. *Nature Physics*, 7(7):565–570, 2011.
- [107] E. Fukushima. *Experimental Pulse NMR: A Nuts and Bolts Approach*. CRC Press, 2018.
- [108] N. Samkharadze, A. Kumar, M. J. Manfra, L. N. Pfeiffer, K. W. West, and G. A. Csáthy. Integrated electronic transport and thermometry at millikelvin temperatures and in strong magnetic fields. *Review of Scientific Instruments*, 82(5):053902, 2011.
- [109] D. I. Bradley, R. E. George, D. Gunnarsson, R. P. Haley, H. Heikkinen, Yu. A. Pashkin, J. Penttilä, J. R. Prance, M. Prunnila, L. Roschier, and M. Sarsby. Nanoelectronic primary thermometry below 4 mK. *Nature Communications*, 7(1):10455, 2016.
- [110] L. A. De Lorenzo and K. C. Schwab. Ultra-high Q acoustic resonance in superfluid ^4He . *Journal of Low Temperature Physics*, 186(3):233–240, Feb 2017.
- [111] G. I. Harris, D. L. McAuslan, E. Sheridan, Y. Sachkou, C. Baker, and W. P. Bowen. Laser cooling and control of excitations in superfluid helium. *Nature Physics*, 12:788 EP –, 04 2016.
- [112] L. Childress, M. P. Schmidt, A. D. Kashkanova, C. D. Brown, G. I. Harris, A. Aiello, F. Marquardt, and J. G. E. Harris. Cavity optomechanics in a levitated helium drop. *Phys. Rev. A*, 96:063842, Dec 2017.

- [113] A. D. Kashkanova, A. B. Shkarin, C. D. Brown, N. E. Flowers-Jacobs, L. Childress, S. W. Hoch, L. Hohmann, K. Ott, J. Reichel, and J. G. E. Harris. Superfluid brillouin optomechanics. *Nature Physics*, 13(1):74–79, 2017.
- [114] P.M. Platzman and M.I. Dykman. Quantum computing with electrons floating on liquid helium. *Science*, 284:1967, 1999.
- [115] M. I. Dykman, P. M. Platzman, and P. Seddighrad. Qubits with electrons on liquid helium. *Phys. Rev. B*, 67:155402, Apr 2003.
- [116] S.A. Lyon. Spin-based quantum computing using electrons on liquid helium. *Physical Review A*, 74:052338, 2006.
- [117] D. I. Schuster, A. Fragner, M. I. Dykman, S. A. Lyon, and R. J. Schoelkopf. Proposal for manipulating and detecting spin and orbital states of trapped electrons on helium using cavity quantum electrodynamics. *Phys. Rev. Lett.*, 105:040503, Jul 2010.
- [118] Fei Yan, Simon Gustavsson, Archana Kamal, Jeffrey Birenbaum, Adam P Sears, David Hover, Ted J. Gudmundsen, Danna Rosenberg, Gabriel Samach, S Weber, Jonilyn L. Yoder, Terry P. Orlando, John Clarke, Andrew J. Kerman, and William D. Oliver. The flux qubit revisited to enhance coherence and reproducibility. *Nature Communications*, 7:12964 EP –, 11 2016.
- [119] Gerwin Koolstra, Ge Yang, and David I. Schuster. Coupling a single electron on superfluid helium to a superconducting resonator. *Nature Communications*, 10(1):5323, 2019.
- [120] J. R. Lane, D. Tan, N. R. Beysengulov, K. Nasyedkin, E. Brook, L. Zhang, T. Stefanski, H. Byeon, K. W. Murch, and J. Pollanen. Integrating superfluids with superconducting qubit systems. *Phys. Rev. A*, 101:012336, Jan 2020.
- [121] J. F and J.F. Annett. *Superconductivity, Superfluids and Condensates*. Oxford Master Series in Physics. OUP Oxford, 2004.
- [122] Cecil T. Lane. *Superfluid Physics*. McGraw-Hill Book Company, 1962.
- [123] L. Tisza. *C. R. Acad. Sci.*, 207:1035, 1938.
- [124] W. Hartung, J. Bierwagen, S. Bricker, C. Compton, T. Grimm, M. Johnson, D. Meidlinger, D. Pendell, J. Popielarski, L. Saxton, and R. York. RF performand of a superconducting S-band cavity filled with liquid helium. In *Proceedings of LINAC 2006*, 2006.

- [125] N. R. Beysengulov, J. R. Lane, J. M. Kitzman, K. Nasyedkin, D. G. Rees, and J. Pollanen. Noise performance and thermalization of single electron transistors using quantum fluids, 2020.
- [126] R. K. Wangsness and F. Bloch. The dynamical theory of nuclear induction. *Phys. Rev.*, 89:728–739, Feb 1953.
- [127] A. G. Redfield. On the theory of relaxation processes. *IBM Journal of Research and Development*, 1(1):19–31, Jan 1957.
- [128] Lara Faoro and Lev B. Ioffe. Microscopic origin of low-frequency flux noise in Josephson circuits. *Phys. Rev. Lett.*, 100:227005, Jun 2008.
- [129] Oliver Dial, Douglas T McClure, Stefano Poletto, G A Keefe, Mary Beth Rothwell, Jay M Gambetta, David W Abraham, Jerry M Chow, and Matthias Steffen. Bulk and surface loss in superconducting transmon qubits. *Superconductor Science and Technology*, 29(4):044001, 2016.
- [130] J. Wenner, R. Barends, R. C. Bialczak, Yu Chen, J. Kelly, Erik Lucero, Matteo Mariantoni, A. Megrant, P. J. J. O’Malley, D. Sank, A. Vainsencher, H. Wang, T. C. White, Y. Yin, J. Zhao, A. N. Cleland, and John M. Martinis. Surface loss simulations of superconducting coplanar waveguide resonators. *Applied Physics Letters*, 99(11):113513, 2011.
- [131] A. Bruno, G. de Lange, S. Asaad, K. L. van der Enden, N. K. Langford, and L. DiCarlo. Reducing intrinsic loss in superconducting resonators by surface treatment and deep etching of silicon substrates. *Applied Physics Letters*, 106(18):182601, 2015.
- [132] A. K. B. Engebretson and B. Golding. Surface acoustic wave propagation in nb:h at low temperatures. In *1997 IEEE Ultrasonics Symposium Proceedings. An International Symposium (Cat. No.97CH36118)*, volume 1, pages 499–502 vol.1, Oct 1997.
- [133] Grigoriy J. Grabovskij, Torben Peichl, Jürgen Lisenfeld, Georg Weiss, and Alexey V. Ustinov. Strain tuning of individual atomic tunneling systems detected by a superconducting qubit. *Science*, 338(6104):232–234, 2012.
- [134] D. J. Van Harlingen, T. L. Robertson, B. L. T. Plourde, P. A. Reichardt, T. A. Crane, and John Clarke. Decoherence in Josephson-junction qubits due to critical-current fluctuations. *Phys. Rev. B*, 70:064517, Aug 2004.
- [135] P.V. Klimov, J. Kelly, Z. Chen, M. Neeley, A. Megrant, B. Burkett, R. Barends, K. Arya, B. Chiaro, Yu Chen, A. Dunsworth, A. Fowler, B. Foxen, C. Gidney, M. Giustina, R. Graff, T. Huang, E. Jeffrey, E. Lucero, J.Y. Mutus, O. Naaman, C. Neill, C. Quintana, P. Roushan, Daniel Sank, A. Vainsencher, J. Wenner, T. C. White, S. Boixo, R. Babbush, V.N. Smelyanskiy, H. Neven, , and John M. Martinis.

- Fluctuations of energy-relaxation times in superconducting qubits. *Phys. Rev. Lett.*, 121:090502, 2018.
- [136] Clemens Müller, Jürgen Lisenfeld, Alexander Shnirman, and Stefano Poletto. Interacting two-level defects as sources of fluctuating high-frequency noise in superconducting circuits. *Phys. Rev. B*, 92:035442, Jul 2015.
 - [137] Jonathan J. Burnett, Andreas Bengtsson, Marco Scigliuzzo, David Niepce, Marina Kudra, Per Delsing, and Jonas Bylander. Decoherence benchmarking of superconducting qubits. *npj Quantum Information*, 5(1):54, 2019.
 - [138] Steffen Schlör, Jürgen Lisenfeld, Clemens Müller, Alexander Bilmes, Andre Schneider, David P. Pappas, Alexey V. Ustinov, and Martin Weides. Correlating decoherence in transmon qubits: Low frequency noise by single fluctuators. *Phys. Rev. Lett.*, 123:190502, Nov 2019.
 - [139] John M. Martinis, M. Ansmann, and J. Aumentado. Energy decay in superconducting Josephson-junction qubits from nonequilibrium quasiparticle excitations. *Phys. Rev. Lett.*, 103:097002, Aug 2009.
 - [140] D. Ristè, C. C. Bultink, M. J. Tiggelman, R. N. Schouten, K. W. Lehnert, and L. DiCarlo. Millisecond charge-parity fluctuations and induced decoherence in a superconducting transmon qubit. *Nature Communications*, 4(1):1913, 2013.
 - [141] J. Wenner, Yi Yin, Erik Lucero, R. Barends, Yu Chen, B. Chiaro, J. Kelly, M. Lenander, Matteo Mariantoni, A. Megrant, C. Neill, P. J. J. O’Malley, D. Sank, A. Vainsencher, H. Wang, T. C. White, A. N. Cleland, and John M. Martinis. Excitation of superconducting qubits from hot nonequilibrium quasiparticles. *Phys. Rev. Lett.*, 110:150502, Apr 2013.
 - [142] K. Serniak, M. Hays, G. de Lange, S. Diamond, S. Shankar, L. D. Burkhardt, L. Frunzio, M. Houzet, and M. H. Devoret. Hot nonequilibrium quasiparticles in transmon qubits. *Phys. Rev. Lett.*, 121:157701, Oct 2018.
 - [143] Antti P. Vepsäläinen, Amir H. Karamlou, John L. Orrell, Akshunna S. Dogra, Ben Loer, Francisca Vasconcelos, David K. Kim, Alexander J. Melville, Bethany M. Niedzielski, Jonilyn L. Yoder, Simon Gustavsson, Joseph A. Formaggio, Brent A. VanDevender, and William D. Oliver. Impact of ionizing radiation on superconducting qubit coherence. *Nature*, 584(7822):551–556, 2020.
 - [144] Laura Cardani, Francesco Valenti, Nicola Casali, Gianluigi Catelani, Thibault Charpentier, Massimiliano Clemenza, Ivan Colantoni, Angelo Cruciani, Luca Gironi, Lukas Grünhaupt, Daria Gusenkova, Fabio Henriques, Marc Lagoin, Maria Martinez, Giorgio Pettinari, Claudia Rusconi, Oliver Sander, Alexey V. Ustinov, Marc Weber, Wolfgang

- Wernsdorfer, Marco Vignati, Stefano Pirro, and Ioan M. Pop. Reducing the impact of radioactivity on quantum circuits in a deep-underground facility, 2020.
- [145] U. Patel, Ivan V. Pechenezhskiy, B. L. T. Plourde, M. G. Vavilov, and R. McDermott. Phonon-mediated quasiparticle poisoning of superconducting microwave resonators. *Phys. Rev. B*, 96:220501, Dec 2017.
 - [146] A.A. Fragner. *Circuit Quantum Electrodynamics with Electrons on Helium*. PhD thesis, Yale University, 2013.
 - [147] K. Nasyedkin, H. Byeon, L. Zhang, N.R. Beysengulov, J. Milem, S. Hemmerle, R. Loloee, and J. Pollanen. Unconventional field effect transistor composed of electrons floating on liquid helium. *Journal of Physics: Condensed Matter*, 30(465501), 2018.
 - [148] F. Souris, H. Christiani, and J. P. Davis. Tuning a 3d microwave cavity via superfluid helium at millikelvin temperatures. *Applied Physics Letters*, 111(17):172601, 2017.
 - [149] James S. Brooks and Russell J. Donnelly. The calculated thermodynamic properties of superfluid helium-4. *Journal of Physical and Chemical Reference Data*, 6(1):51–104, 1977.
 - [150] L. A. De Lorenzo and K. C. Schwab. Superfluid optomechanics: coupling of a superfluid to a superconducting condensate. *New Journal of Physics*, 16(11):113020, nov 2014.
 - [151] C. T. Rogers and R. A. Buhrman. Nature of single-localized-electron states derived from tunneling measurements. *Phys. Rev. Lett.*, 55:859–862, Aug 1985.
 - [152] M. H. Devoret. *Quantum Fluctuations in Electrical Circuits*, chapter 10. Elsevier Science, 1997.
 - [153] G. Catelani, R. J. Schoelkopf, M. H. Devoret, and L. I. Glazman. Relaxation and frequency shifts induced by quasiparticles in superconducting qubits. *Phys. Rev. B*, 84:064517, Aug 2011.
 - [154] X. Y. Jin, A. Kamal, A. P. Sears, T. Gudmundsen, D. Hover, J. Miloshi, R. Slattery, F. Yan, J. Yoder, T. P. Orlando, S. Gustavsson, and W. D. Oliver. Thermal and residual excited-state population in a 3D transmon qubit. *Phys. Rev. Lett.*, 114:240501, Jun 2015.
 - [155] J. Pollanen, H. Choi, J.P. Davis, B.T. Rolfs, and W.P. Halperin. Low temperature thermal resistance for a new design of silver sinter heat exchanger. *J. Phys. Con. Ser.*, 150:012037, 2009.
 - [156] A. A. Clerk and D. Wahyu Utami. Using a qubit to measure photon-number statistics of a driven thermal oscillator. *Phys. Rev. A*, 75:042302, Apr 2007.

- [157] G. Catelani, S. E. Nigg, S. M. Girvin, R. J. Schoelkopf, and L. I. Glazman. Decoherence of superconducting qubits caused by quasiparticle tunneling. *Phys. Rev. B*, 86:184514, Nov 2012.
- [158] K. Geerlings, Z. Leghtas, I. M. Pop, S. Shankar, L. Frunzio, R. J. Schoelkopf, M. Mirrahimi, and M. H. Devoret. Demonstrating a driven reset protocol for a superconducting qubit. *Phys. Rev. Lett.*, 110:120501, Mar 2013.
- [159] Laura De Lorenzo. *Optomechanics with Superfluid Helium-4*. PhD thesis, California Institute of Technology, 2016.
- [160] D. P. Morgan. *Surface Acoustic Wave Devices*. Elsevier Ltd., 2nd edition, 2005.
- [161] Per Delsing, Andrew N Cleland, Martin J A Schuetz, Johannes Knörzer, Géza Giedke, J Ignacio Cirac, Kartik Srinivasan, Marcelo Wu, Krishna Coimbatore Balram, Christopher Bäuerle, Tristan Meunier, Christopher J B Ford, Paulo V Santos, Edgar Cerdas-Méndez, Hailin Wang, Hubert J Krenner, Emeline D S Nysten, Matthias Weiß, Geoff R Nash, Laura Thevenard, Catherine Gourdon, Pauline Rovillain, Max Marangolo, Jean-Yves Duquesne, Gerhard Fischerauer, Werner Ruile, Alexander Reiner, Ben Paschke, Dmytro Denysenko, Dirk Volkmer, Achim Wixforth, Henrik Bruus, Martin Wiklund, Julien Reboud, Jonathan M Cooper, YongQing Fu, Manuel S Brugger, Florian Rehfeldt, and Christoph Westerhausen. The 2019 surface acoustic waves roadmap. *Journal of Physics D: Applied Physics*, 52(35):353001, jul 2019.
- [162] Amir H. Safavi-Naeini and Oskar Painter. *Optomechanical Crystal Devices*, pages 195–231. Springer Berlin Heidelberg, Berlin, Heidelberg, 2014.
- [163] Lord Rayleigh. On Waves Propagated along the Plane Surface of an Elastic Solid. *Proceedings of the London Mathematical Society*, s1-17(1):4–11, 11 1885.
- [164] L. D. Landau and E. M. Lifshitz. *Theory of Elasticity*, volume 7 of Course of Theoretical Physics. Pergamon Press, 2nd edition, 1970.
- [165] K. J. Satzinger. *Quantum control of surface acoustic wave phonons*. PhD thesis, UC Santa Barbara, 2018.
- [166] H. Byeon. *Studying electrons on helium via surface acoustic wave techniques*. PhD thesis, Michigan State University, 2021.
- [167] Göran Johansson and R. H. Hadfield. *Superconducting Devices in Quantum Optics*. Springer, 2016.
- [168] L.A. Tracy. *Studies of two dimensional electron systems via surface acoustic waves and nuclear magnetic resonance techniques*. PhD thesis, California Institute of Technology, Sep 2007.

- [169] R. Manenti, M. J. Peterer, A. Nersisyan, E. B. Magnusson, A. Patterson, and P. J. Leek. Surface acoustic wave resonators in the quantum regime. *Phys. Rev. B*, 93:041411, Jan 2016.
- [170] J. R. Lane. Coupling of modes analysis for saw devices. https://github.com/jrlane2/SAW_COM.
- [171] K.S. Van Dyke. The piezo-electric resonator and its equivalent network. *Proceedings of the Institute of Radio Engineers*, 16(6):742–764, 1928.
- [172] A. Bienfait, K. J. Satzinger, Y. P. Zhong, H.-S. Chang, M.-H. Chou, C. R. Conner, É. Dumur, J. Grebel, G. A. Peairs, R. G. Povey, and A. N. Cleland. Phonon-mediated quantum state transfer and remote qubit entanglement. *Science*, 364(6438):368–371, 2019.
- [173] Andreas Ask, Maria Ekström, Per Delsing, and Göran Johansson. Cavity-free vacuum-rabi splitting in circuit quantum acoustodynamics. *Phys. Rev. A*, 99:013840, Jan 2019.
- [174] M. K. Ekström, T. Aref, A. Ask, G. Andersson, B. Suri, H. Sanada, G. Johansson, and P. Delsing. Towards phonon routing: controlling propagating acoustic waves in the quantum regime. *New Journal of Physics*, 21(12):123013, dec 2019.
- [175] Vitaly S. Shumeiko. Quantum acousto-optic transducer for superconducting qubits. *Phys. Rev. A*, 93:023838, Feb 2016.
- [176] R. S. Weis and T. K. Gaylord. Lithium niobate: Summary of physical properties and crystal structure. *Applied Physics A*, 37(4):191–203, 1985.
- [177] D.L.T. Bell and R.C.M. Li. Surface-acoustic-wave resonators. *Proceedings of the IEEE*, 64(5):711–721, 1976.
- [178] K. J. Satzinger, C. R. Conner, A. Bienfait, H.-S. Chang, Ming-Han Chou, A. Y. Cleland, É. Dumur, J. Grebel, G. A. Peairs, R. G. Povey, S. J. Whiteley, Y. P. Zhong, D. D. Awschalom, D. I. Schuster, and A. N. Cleland. Simple non-galvanic flip-chip integration method for hybrid quantum systems. *Applied Physics Letters*, 114(17):173501, 2019.
- [179] Marco Scigliuzzo, Laure E Bruhat, Andreas Bengtsson, Jonathan J Burnett, Anita Fardavi Roudsari, and Per Delsing. Phononic loss in superconducting resonators on piezo-electric substrates. *New Journal of Physics*, 22(5):053027, may 2020.
- [180] Thomas Pastureaud, Marc Solal, Beatrice Biasse, Bernard Aspar, Jean-bernard Briot, William Daniau, William Steichen, Raphael Lardat, Vincent Laude, Alain Laens, Jean-michel Friedt, and Sylvain Ballandras. High-frequency surface acoustic waves excited

- on thin-oriented linbo3 single-crystal layers transferred onto silicon. *IEEE Transactions on Ultrasonics, Ferroelectrics, and Frequency Control*, 54(4):870–876, 2007.
- [181] M. D. Reed, B. R. Johnson, A. A. Houck, L. DiCarlo, J. M. Chow, D. I. Schuster, L. Frunzio, and R. J. Schoelkopf. Fast reset and suppressing spontaneous emission of a superconducting qubit. *Applied Physics Letters*, 96(20):203110, 2010.
 - [182] K. W. Murch, U. Vool, D. Zhou, S. J. Weber, S. M. Girvin, and I. Siddiqi. Cavity-assisted quantum bath engineering. *Phys. Rev. Lett.*, 109:183602, Oct 2012.
 - [183] M. Naghiloo, M. Abbasi, Yogesh N. Joglekar, and K. W. Murch. Quantum state tomography across the exceptional point in a single dissipative qubit. *Nature Physics*, 15(12):1232–1236, 2019.
 - [184] R. L. Willett, R. R. Ruel, M. A. Pallanan, K. W. West, and L.N. Pfeiffer. Enhanced finite-wave-vector conductivity at multiple even-denominator filling factors in two-dimensional electron systems. *Phys. Rev. B*, 47(12):7344, 1993.
 - [185] R. L. Willett, R. R. Ruel, K. W. West, and L. N. Pfeiffer. Experimental demonstration of a Fermi surface at one-half filling of the lowest Landau level. *Phys. Rev. Lett.*, 71:3846–3849, Dec 1993.
 - [186] David Tong. Lectures on the quantum Hall effect. arXiv preprint: 1606.06687, Jun 2016.
 - [187] J. R. Lane, L. Zhang, M. A. Khasawneh, B. N. Zhou, E. A. Henriksen, and J. Pollanen. Flip-chip gate-tunable acoustoelectric effect in graphene. *Journal of Applied Physics*, 124(19):194302, 2018.
 - [188] P. R. Wallace. The band theory of graphite. *Phys. Rev.*, 71:622–634, May 1947.
 - [189] K. S. Novoselov, A. K. Geim, S. V. Morozov, D. Jiang, Y. Zhang, S. V. Dubonos, I. V. Grigorieva, and A. A. Firsov. Electric field effect in atomically thin carbon films. *Science*, 306(5696):666–669, 2004.
 - [190] J. I. A. Li, Q. Shi, Y. Zeng, K. Watanabe, T. Taniguchi, J. Hone, and C. R. Dean. Pairing states of composite fermions in double-layer graphene. *Nature Physics*, 15(9):898–903, 2019.
 - [191] K. S. Novoselov, A. K. Geim, S. V. Morozov, D. Jiang, M. I. Katsnelson, I. V. Grigorieva, S. V. Dubonos, and A. A. Firsov. Two-dimensional gas of massless Dirac fermions in graphene. *Nature*, 438(7065):197–200, 2005.

- [192] Yuanbo Zhang, Yan-Wen Tan, Horst L. Stormer, and Philip Kim. Experimental observation of the quantum Hall effect and Berry's phase in graphene. *Nature*, 438(7065):201–204, 2005.
- [193] L. Wang, I. Meric, P. Y. Huang, Q. Gao, Y. Gao, H. Tran, T. Taniguchi, K. Watanabe, L. M. Campos, D. A. Muller, J. Guo, P. Kim, J. Hone, K. L. Shepard, and C. R. Dean. One-dimensional electrical contact to a two-dimensional material. *Science*, 342(6158):614–617, 2013.
- [194] Xu Du, Ivan Skachko, Fabian Duerr, Adina Luican, and Eva Y. Andrei. Fractional quantum Hall effect and insulating phase of Dirac electrons in graphene. *Nature*, 462(7270):192–195, 2009.
- [195] Kirill I. Bolotin, Fereshte Ghahari, Michael D. Shulman, Horst L. Stormer, and Philip Kim. Observation of the fractional quantum Hall effect in graphene. *Nature*, 462(7270):196–199, 2009.
- [196] J. I. A. Li, C. Tan, S. Chen, Y. Zeng, T. Taniguchi, K. Watanabe, J. Hone, and C. R. Dean. Even-denominator fractional quantum Hall states in bilayer graphene. *Science*, 358(6363):648–652, 2017.
- [197] B. Hunt, J. D. Sanchez-Yamagishi, A. F. Young, M. Yankowitz, B. J. LeRoy, K. Watanabe, T. Taniguchi, P. Moon, M. Koshino, P. Jarillo-Herrero, and R. C. Ashoori. Massive Dirac fermions and Hofstadter butterfly in a van der Waals heterostructure. *Science*, 340(6139):1427–1430, 2013.
- [198] Yuan Cao, Valla Fatemi, Shiang Fang, Kenji Watanabe, Takashi Taniguchi, Efthimios Kaxiras, and Pablo Jarillo-Herrero. Unconventional superconductivity in magic-angle graphene superlattices. *Nature*, 556(7699):43–50, 2018.
- [199] C. W. J. Beenakker. Colloquium: Andreev reflection and Klein tunneling in graphene. *Rev. Mod. Phys.*, 80:1337–1354, Oct 2008.
- [200] K. S. Novoselov, Z. Jiang, Y. Zhang, S. V. Morozov, H. L. Stormer, U. Zeitler, J. C. Maan, G. S. Boebinger, P. Kim, and A. K. Geim. Room-temperature quantum Hall effect in graphene. *Science*, 315(5817):1379–1379, 2007.
- [201] A. R. Hutson and Donald L. White. Elastic wave propagation in piezoelectric semiconductors. *Journal of Applied Physics*, 33(1):40–47, 1962.
- [202] A. Wixforth, J. Scriba, M. Wassermeier, J. P. Kotthaus, G. Weimann, and W. Schlapp. Surface acoustic waves on GaAs/Al_xGa_{1-x}As heterostructures. *Phys. Rev. B*, 40(11):7874, 1989.

- [203] A. Wixforth, J. P. Kotthaus, and G. Weimann. Quantum oscillations in the surface-acoustic-wave attenuation caused by a two-dimensional electron system. *Phys. Rev. Lett.*, 56:2104, 1986.
- [204] A. Esslinger, A. Wixforth, R. W. Winkler, J. P. Kotthaus, H. Nickel, W. Schlapp, and R. Lösch. Acoustoelectric study of localized states in the quantized Hall effect. *Solid State Communications*, 84(10):939, 1992.
- [205] I. V. Kukushkin, V. Umansky, K. von Klitzing, and J. H. Smet. Collective modes and the periodicity of quantum Hall stripes. *Phys. Rev. Lett.*, 106:206804, May 2011.
- [206] J. Pollanen, J.P. Eisenstein, L.N. Pfeiffer, and K.W. West. Charge metastability and hysteresis in the quantum Hall regime. *Phys. Rev. B*, 94(24):245440, 2016.
- [207] B. Friess, Y. Peng, B. Rosenow, F. von Oppen, V. Umansky, K. von Klitzing, and J. H. Smet. Negative permittivity in bubble and stripe phases. *Nat Phys*, doi:10.1038/nphys4213, 2017.
- [208] M. A. Paalanen, R. L. Willett, P. B. Littlewood, R. R. Ruel, K. W. West, L. N. Pfeiffer, and D. J. Bishop. rf conductivity of a two-dimensional electron system at small Landau-level filling factors. *Phys. Rev. B*, 45:11342–11345, May 1992.
- [209] I. L. Drichko, I. Y. Smirnov, A. V. Suslov, L. N. Pfeiffer, K. W. West, and Y. M. Galperin. Crossover between localized states and pinned Wigner crystal in high-mobility n-GaAs/AlGaAs heterostructures near filling factor $\nu = 1$. *Phys. Rev. B*, 92:205313, 2015.
- [210] A. V. Suslov, I. L. Drichko, I. Y. Smirnov, A. F. Ioffe, L. N. Pfeiffer, K. W. West, and Y. M. Galperin. Study of Wigner crystal in n-GaAs/AlGaAs by surface acoustic waves. *The Journal of the Acoustical Society of America*, 138:1938, 2015.
- [211] L.A. Tracy, J.P. Eisenstein, M.P. Lilly, L.N. Pfeiffer, and K.W. West. Surface acoustic wave propagation and inhomogeneities in low-density two-dimensional electron systems near the metal–insulator transition. *Solid State Communications*, 137(3):150 – 153, 2006.
- [212] A. L. Efros and Yu. M. Galperin. Quantization of the acoustoelectric current in a two-dimensional electron system in a strong magnetic field. *Phys. Rev. Lett.*, 64:1959–1962, Apr 1990.
- [213] A. Esslinger, R.W. Winkler, C. Rocke, A. Wixforth, J.P. Kotthaus, H. Nickel, W. Schlapp, and R. Lösch. Ultrasonic approach to the integer and fractional quantum Hall effect. *Surface Science*, 305(1):83 – 86, 1994.

- [214] J. M. Shilton, D. R. Mace, V. I. Talyanskii, M. Y. Simmons, M. Pepper, A. C. Churchill, and D. A. Ritchie. Experimental study of the acoustoelectric effects in GaAs-AlGaAs heterostructures. *J. Phys. Condens. Matter*, 7:7675, 1995.
- [215] M. Rotter, A. Wixforth, W. Ruile, D. Bernklau, and H. Riechert. Giant acoustoelectric effect in GaAs/LiNbO₃ hybrids. *App. Phys. Lett.*, 73(15):2128, 1998.
- [216] J. M. Shilton, D. R. Mace, V. I. Talyanskii, M. Pepper, M. Y. Simmons, A. C. Churchill, and D. A. Ritchie. Effect of spatial dispersion on acoustoelectric current in a high-mobility two-dimensional electron gas. *Phys. Rev. B*, 51:14770–14773, May 1995.
- [217] J. Ebbecke, N. E. Fletcher, T. J. B. M. Janssen, H. E. Beere, D. A. Ritchie, and M. Pepper. Acoustoelectric current transport through a double quantum dot. *Phys. Rev. B*, 72(121311(R)), 2005.
- [218] M. Kataoka, M. R. Astley, A. L. Thorn, D. K. L. Oi, C. H. W. Barnes, C. J. B. Ford, D. Anderson, G. A. C. Jones, I. Farrer, D. A. Ritchie, and M. Pepper. Coherent time evolution of a single-electron wave function. *Phys. Rev. Lett.*, 102:156801, 2009.
- [219] J. M. Shilton, D. R. Mace, V. I. Talyanskii, Yu Galperin, M. Y. Simmons, M. Pepper, and D. A. Ritchie. On the acoustoelectric current in a one-dimensional channel. *J. Phys. Condens. Matter*, 8:L337, 1996.
- [220] J. M. Shilton, V. I. Talyanskii, M. Pepper, D. A. Ritchie, J. E. F. Frost, C. J. B. Ford, C. G. Smith, and G. A. C. Jones. High-frequency single electron transport in a quasi-one-dimensional GaAs channel induced by surface acoustic waves. *J. Phys. Condens. Matter*, 8:L531, 1996.
- [221] M. R. Astley, M. Kataoka, R. J. Schneble, C. J. B. Ford, C. H. W. Barnes, D. Anderson, G. A. C. Jones, H. E. Beere, D. A. Ritchie, and M. Pepper. Examination of the surface acoustic wave reflections by observing acoustoelectric current generation under pulse modulation. *Applied Physics Letters*, 89:132102, 2006.
- [222] H. Byeon, K. Nasyedkin, J. R. Lane, N. R. Beysengulov, L. Zhang, R. Loloee, and J. Pollanen. Piezoacoustics for flying electron qubits on helium. arXiv:2008.02330, Aug 2020.
- [223] Peter Thalmeier, Balázs Dóra, and Klaus Ziegler. Surface acoustic wave propagation in graphene. *Phys. Rev. B*, 81:041409, Jan 2010.
- [224] S. H. Zhang and W. Xu. Absorption of surface acoustic waves by graphene. *AIP Advances*, 1(2):022146, 2011.

- [225] Jürgen Schiefele, Jorge Pedrós, Fernando Sols, Fernando Calle, and Francisco Guinea. Coupling light into graphene plasmons through surface acoustic waves. *Phys. Rev. Lett.*, 111:237405, Dec 2013.
- [226] Z. Insepov, E. Emelin, O. Kononenko, D. V. Roshchupkin, K. B. Tnyshtykbayev, and K. A. Baigarin. Surface acoustic wave amplification by direct current-voltage supplied to graphene film. *Applied Physics Letters*, 106(2):023505, 2015.
- [227] V. Miseikis, J. E. Cunningham, K. Saeed, R. O’Rourke, and A. G. Davies. Acoustically induced current flow in graphene. *App. Phys. Let.*, 100:133105, 2012.
- [228] V. Miseikis. *The Interaction of Graphene with High-Frequency Acoustic and Electromagnetic Waves*. PhD thesis, University of Leeds, 2012.
- [229] L. Bandhu, L. M. Lawton, and G. R. Nash. Macroscopic acoustoelectric charge transport in graphene. *Applied Physics Letters*, 103(13):133101, 2013.
- [230] P. V. Santos, T. Schumann, M. H. Oliveira Jr., J. M. J. Lopes, and H. Riechert. Acousto-electric transport in epitaxial monolayer graphene on sic. *Applied Physics Letters*, 102(22):221907, 2013.
- [231] Lokeshwar Bandhu and Geoffrey R. Nash. Controlling the properties of surface acoustic waves using graphene. *Nano Research*, 9(3):685–691, 2016.
- [232] T. Poole and G. R. Nash. Acoustoelectric current in graphene nanoribbons. *Scientific Reports*, 7(1):1767, 2017.
- [233] A. Schenstrom, Y.J. Qian, M.-F. Xu, H.-P. Baum, M. Levy, and Bimal K. Sarma. Oscillations in the acousto-electric proximity coupling to a 2d electron gas. *Solid State Communications*, 65(7):739 – 742, 1988.
- [234] G. S. Kino and T. M. Reeder. A normal mode theory for the rayleigh wave amplifier. *IEEE Transactions on Electronic Devices*, 18(10):909–920, 1971.
- [235] A. Wixforth, J. Scriba, M. Wassermeier, and J. P. Kotthaus. Interaction of surface acoustic waves with a two-dimensional electron system in a LiNbO₃-GaAs/AlGaAs sandwich structure. *J. Appl. Phys.*, 64(4):2213, 1988.
- [236] Bennaceur Keyan, Benjamin A. Schmidt, Samuel Gaucher, Dominique Laroche, Michael P. Lilly, John L. Reno, Ken W. West, Loren N. Pfeiffer, and Guillaume Gervais. Mechanical flip-chip for ultra-high electron mobility devices. *Scientific Reports*, 5:13494 EP –, 09 2015.

- [237] J. Martin, N. Akerman, G. Ulbricht, T. Lohmann, J. H. Smet, K. von Klitzing, and A. Yacoby. Observation of electron-hole puddles in graphene using a scanning single-electron transistor. *Nat Phys*, 4(2):144–148, 02 2008.
- [238] M.-Y. Li, C.-C. Tang, D.C. Ling, C.C. Chi, and J.-C. Chen. Charged impurity-induced scatterings in chemical vapor deposited graphene. *J. Appl. Phys.*, 114:233703, 2013.
- [239] S. Okuda, T. Ikuta, Y. Kanai, T. Ono, S. Ogawa, D. Fujisawa, M. Shimatani, K. Inoue, K. Maehashi, and K. Matsumoto. Acoustic carrier transportation induced by surface acoustic waves in graphene in solution. *Applied Physics Express*, 9:045104, 2016.
- [240] C. Tang, Y. Chen, D.C. Ling, C.C. Chi, and J. Chen. Ultra-low acoustoelectric attenuation in graphene. *J. Appl. Phys.*, 121:124505, 2017.
- [241] J Wenner, M Neeley, Radoslaw C Bialczak, M Lenander, Erik Lucero, A D O’Connell, D Sank, H Wang, M Weides, A N Cleland, and John M Martinis. Wirebond crosstalk and cavity modes in large chip mounts for superconducting qubits. *Superconductor Science and Technology*, 24(6):065001, 2011.
- [242] Md. Shafayat Hossain, Meng K. Ma, M. A. Mueed, L. N. Pfeiffer, K. W. West, K. W. Baldwin, and M. Shayegan. Direct observation of composite fermions and their fully-spin-polarized fermi sea near $\nu = 5/2$. *Phys. Rev. Lett.*, 120:256601, Jun 2018.
- [243] Gustav Andersson, Maria K. Ekström, and Per Delsing. Electromagnetically induced acoustic transparency with a superconducting circuit. *Phys. Rev. Lett.*, 124:240402, Jun 2020.
- [244] B. G. Christensen, C. D. Wilen, A. Opremcek, J. Nelson, F. Schlenker, C. H. Zimonick, L. Faoro, L. B. Ioffe, Y. J. Rosen, J. L. DuBois, B. L. T. Plourde, and R. McDermott. Anomalous charge noise in superconducting qubits. *Phys. Rev. B*, 100:140503, Oct 2019.
- [245] Fei Yan, Simon Gustavsson, Jonas Bylander, Xiaoyue Jin, Fumiki Yoshihara, David G. Cory, Yasunobu Nakamura, Terry P. Orlando, and William D. Oliver. Rotating-frame relaxation as a noise spectrum analyser of a superconducting qubit undergoing driven evolution. *Nature Communications*, 4:2337 EP –, 08 2013.